

К. П. Яковлев

**МАТЕМАТИЧЕСКАЯ
ОБРАБОТКА
РЕЗУЛЬТАТОВ
ИЗМЕРЕНИЙ**



К. П. ЯКОВЛЕВ

**МАТЕМАТИЧЕСКАЯ
ОБРАБОТКА
РЕЗУЛЬТАТОВ
ИЗМЕРЕНИЙ**

ИЗДАНИЕ ВТОРОЕ
ИСПРАВЛЕННОЕ

ГОСУДАРСТВЕННОЕ ИЗДАТЕЛЬСТВО
ТЕХНИКО-ТЕОРЕТИЧЕСКОЙ ЛИТЕРАТУРЫ
МОСКВА 1953

К. П. ЯКОВЛЕВ

**МАТЕМАТИЧЕСКАЯ
ОБРАБОТКА
РЕЗУЛЬТАТОВ
ИЗМЕРЕНИЙ**

**ИЗДАНИЕ ВТОРОЕ
ИСПРАВЛЕННОЕ**

**ГОСУДАРСТВЕННОЕ ИЗДАТЕЛЬСТВО
ТЕХНИКО-ТЕОРЕТИЧЕСКОЙ ЛИТЕРАТУРЫ
МОСКВА 1953**

АННОТАЦИЯ

Книга посвящена вопросам математической обработки результатов измерений, большое внимание уделено методической стороне. Изложение материала иллюстрируется достаточным количеством интересных примеров.

Общий характер книги сохранён и во втором издании, в котором были исправлены только опечатки и ошибки первого издания, а также внесены небольшие изменения текста в тех местах, где автор считал изложение недостаточно точным.

Книга доступна для студентов физических и физико-математических факультетов; она может быть использована как пособие работающим в области физики и физической химии, а также инженерами и техниками, ведущими экспериментальную работу.

ОПЕЧАТКИ

Стр.	Строка	Напечатано	Должно быть
44	12 св.	$6,62 \cdot 10^{27}$	$6,62 \cdot 10^{-27}$
44	7 сн.	4,8029	$4,8029 \cdot 10^{-10}$
189	3 св.	n	n^2
196	12 сн.	$\sum $	$\sum \varepsilon $
245	форм. (6)	y_{1i}	y_1
		y_{2i}	y_2
		y_{3i}	y_3
		$\overline{y_{1i}}$	$\overline{y_1}$
246	форм. (6')	$\overline{y_{2i}}$	$\overline{y_2}$
		$\overline{y_{3i}}$	$\overline{y_3}$
		они с буквой Z	с буквой Z они
251	13 сн.		a_{11}
323	12 сн.	a_{10}	b_{11}
323	9 сн.	b_{10}	$\beta \pm$
373	9 сн.	$\beta \mp$	$\beta \pm$
373	1 сн.	$(\varphi \mp \alpha)$	$(\varphi \pm \alpha)$

Зак. 1283

3. Определение наиболее выгодных условий измерения	120
Глава IV. Закон нормального распределения случайных ошибок	128
1. Случайные явления и их общая классификация	128
2. Математическое определение вероятности	131

АННОТАЦИЯ

Книга посвящена вопросам математической обработки результатов измерений, большое внимание уделено методической стороне. Изложение материала иллюстрируется достаточным количеством интересных примеров.

Общий характер книги сохранён и во втором издании, в котором были исправлены только опечатки и ошибки первого издания, а также внесены небольшие изменения текста в тех местах, где автор считал изложение недостаточно точным.

Книга доступна для студентов физических и физико-математических факультетов; она может быть использована как пособие работающим в области физики и физической химии, а также инженерами и техниками, ведущими экспериментальную работу.

ОГЛАВЛЕНИЕ

От автора	6
Предисловие к первому изданию	7
В в е д е н и е. Основные вопросы математической обработки результатов измерений	9
Г л а в а I. Приближённые значения величин и их ошибки	18
1. Приближённые значения величин	18
2. Округление приближённых значений чисел; правило дополнения	19
3. Ошибки приближённых значений	21
4. Знак ошибок приближённых значений	26
5. Классификация приближённых чисел	29
6. Верные цифры в приближённых значениях чисел	33
Г л а в а II. Основные арифметические действия с приближёнными значениями чисел	47
1. Малые величины различных порядков	47
2. Формулы для приближённых вычислений	51
3. Сложение и вычитание приближённых чисел	57
4. Умножение и деление приближённых чисел	66
5. Возведение в степень приближённых чисел и извлечение из них корня	72
6. Несколько практических указаний	76
Г л а в а III. Ошибки приближённых значений функций и общая теория ошибок	80
1. Основные задачи теории ошибок	80
2. Ошибки функции одного и двух независимых переменных	85
3. Ошибки функции нескольких независимых переменных	98
4. Вторая задача теории ошибок	108
5. Определение наиболее выгодных условий измерения	120
Г л а в а IV. Закон нормального распределения случайных ошибок	128
1. Случайные явления и их общая классификация	128
2. Математическое определение вероятности	131

3. Теоремы сложения и умножения вероятностей . . .	134
4. Общие закономерности случайных ошибок . . .	139
5. Основная формула теории случайных ошибок . . .	142
6. Интеграл вероятностей и его вычисление . . .	151
Глава V. Показатели точности измерений . . .	159
1. Постулат среднего арифметического . . .	159
2. Мера точности для точных значений ошибок и для отклонений от среднего арифметического . . .	162
3. Ошибки измерений—средняя квадратичная, вероятная и средняя . . .	171
4. Геометрическое значение средней квадратичной, вероятной и средней ошибок отдельных измерений . . .	181
5. Ошибки среднего арифметического . . .	184
6. Взвешенные измерения и исправление результатов . . .	190
Глава VI. Основные приёмы графического анализа результатов измерений . . .	201
1. Графическое изображение результатов измерений . . .	201
2. Графическое дифференцирование и интегрирование . . .	211
3. Функциональные шкалы и их применение . . .	219
4. Логарифмическая шкала, степенная шкала и проективная шкала . . .	226
5. Координатные сетки . . .	231
6. Спрямление кривых при помощи функциональных сеток . . .	232
Глава VII. Элементы номографии . . .	234
1. Общее понятие о номографии . . .	234
2. Сетчатые номограммы . . .	236
3. Номограммы на выравненных точках . . .	245
4. Зет-номограммы . . .	250
5. Примеры расчёта и построения номограмм . . .	252
Глава VIII. Основные приёмы интерполирования . . .	257
1. Ливейное и графическое интерполирование . . .	257
2. Интерполирование при равноотстоящих значениях аргумента . . .	260
3. Формула Ньютона . . .	271
4. Интерполирование при неравноотстоящих значениях аргумента . . .	285
5. Точность интерполяционных формул . . .	289
Глава IX. Основы гармонического анализа . . .	294
1. Периодические процессы . . .	294
2. Гармонический анализ . . .	297
3. Схема вычислений для двенадцати ординат . . .	303
4. Примеры вычисления коэффициентов разложения при двенадцати ординатах . . .	312
5. Технические приёмы гармонического анализа . . .	320

Глава X. Эмпирические формулы . . .	326
1. Общие замечания . . .	326
2. Графический метод . . .	331
3. Метод интерполяционных формул . . .	342
4. Метод средних . . .	349
5. Метод наименьших квадратов . . .	358
6. Выбор формул лучшего типа . . .	367
Приложения . . .	373
1. Формулы для приближённых вычислений . . .	373
2. Абсолютные и относительные ошибки некоторых функций . . .	374
3. Значения интеграла вероятностей P . . .	275
Рекомендуемая литература . . .	379
Предметный указатель . . .	380

ОТ АВТОРА

Общий характер изложения, принятый в первом издании «Математической обработки результатов измерений», сохранён мною и во втором издании. Все изменения в этом издании ограничиваются, помимо исправления нескольких опечаток и ошибок, оказавшихся в тексте, только переработкой тех мест, где изложение мне казалось недостаточно точным или недостаточно ясным, и несколькими небольшими дополнениями, главным образом в главе V.

Несколько указаний на небольшие неточности, имеющиеся в книге, я получил от некоторых читателей, за что и приношу им свою благодарность. Все дальнейшие указания на недостатки книги и на желательные её изменения и дополнения, которые могут последовать после выхода второго издания, будут приняты мной с не меньшей признательностью.

Сентябрь 1953 г.

Проф. *К. П. Яковлев*

ПРЕДИСЛОВИЕ К ПЕРВОМУ ИЗДАНИЮ

Вопросам приближённых вычислений и численного математического анализа посвящено большое число монографий и руководств. Достаточно указать хотя бы на классический труд акад. А. Н. Крылова «Лекции о приближённых вычислениях», а также на такие переводные книги, как «Численные методы математического анализа» Дж. Скарборо и «Элементы прикладного анализа» Г. Зандена. В противоположность этому литература по вопросам обработки результатов измерений, т. е. по вопросам, близким к задачам приближённых вычислений, но всё же несколько специальным, представлена значительно меньшим числом книг. Здесь можно указать только одну монографию Э. Уиттекера и Г. Робинсона «Математическая обработка результатов наблюдений», монографию, в которой собран полный материал, строго проработанный, но не вполне пригодный для повседневной справочной работы.

Надеясь пополнить хотя бы частично этот недостаток литературы по вопросам обработки результатов измерений, я счёл возможным переработать и подготовить для печати те лекции по курсу «Математическая обработка результатов измерений», которые я в течение нескольких лет читал для студентов физического факультета Московского государственного университета. К этому меня побуждали также большой интерес к моим лекциям и настойчивое желание их прослушать, которое неоднократно выражали научные работники ряда московских исследовательских институтов. Почти весь материал, излагавшийся мной на лекциях, вошёл в настоящую монографию, хотя и с некоторыми сокращениями, неизбежными в печатном издании. Но вместе с тем необходимо отметить, что

эта книга совершенно не претендует на полноту изложения очень разнообразных вопросов, которые составляют её содержание, так же как она не претендует и на математическую строгость доказательств многочисленных положений и теорем, которыми приходится пользоваться: в этой монографии излагается в очень краткой и сжатой форме только тот материал, овладеть которым необходимо всем, кто, занимаясь измерительными работами в лаборатории, одновременно рассчитывает получить на основании результатов своих измерений ряд выводов и обобщений.

Изложение всего материала я старался вести в возможно более простой форме, обращаясь к высшей математике только там, где это существенно необходимо, но и здесь я ограничивался только её основными общеизвестными приёмами. Все положения и формулы после теоретического обоснования рассматриваются затем в их практических применениях на примерах, взятых из научных исследований или из повседневной лабораторной практики. Я считаю безусловно необходимым рассматривать достаточно большое количество примеров, так как такой метод даёт возможность читателям быстро приобретать практические навыки, крайне ценные. Без этих практических навыков не представляется возможным быстро и уверенно находить правильное решение и преодолевать те многочисленные затруднения, с которыми мы неизбежно встречаемся, применяя на практике теоретические положения и формулы и доводя все расчёты до конца, т. е. до окончательного численного результата.

Я не имел в виду придавать этой книге характер справочника, но думаю, что в ней можно найти ответы почти на все вопросы, которые ставит перед нами задача обработки результатов измерений. Возможно, однако, что в этой книге что-либо изложено недостаточно полно или упущено совершенно. Вследствие этого я буду весьма признателен за все указания на недостатки и дефекты книги, а также на желательные её изменения и дополнения.

Июль 1950 г.

Проф. К. П. Яковлев

ВВЕДЕНИЕ

ОСНОВНЫЕ ВОПРОСЫ МАТЕМАТИЧЕСКОЙ ОБРАБОТКИ РЕЗУЛЬТАТОВ ИЗМЕРЕНИЙ

Основной задачей опытной физики принято считать количественное исследование физических явлений, что даёт нам возможность определять числовые значения физических величин и устанавливать в конечном результате законы исследуемых явлений. При этих работах мы неизбежно встречаемся с двумя последовательными процессами: во-первых, с измерениями, которые выполняются в процессе лабораторной работы, и, во-вторых, с вычислениями, которые выполняются по окончании измерений при обработке результатов. Как при измерениях, так и при вычислениях перед нами всегда возникают некоторые затруднения, различные в том и другом случае.

Затруднения при экспериментальных работах определяются главным образом необходимостью обеспечить отчётливое и точное действие аппаратуры, её проверкой, контролем за постоянством её показаний и т. п. Экспериментатору часто приходится встречаться с чрезвычайно большими трудностями и осложнениями в своей работе, в особенности, если применяют сложные измерительные методы, или разрабатывают новую методику измерений, или, наконец, производят измерения с высокой степенью точности. Для того, чтобы в таких условиях успешно закончить измерения, требуются большие знания и навыки, большая опытность во всех вопросах, связанных с постановкой и проведением лабораторных исследований, что обыкновенно приобретает только в результате многолетней практики.

Трудности, которые мы встречаем, приступая к обработке результатов измерений, хорошо известны всем, кому приходилось решать достаточно сложные вычислительные задачи. При выполнении вычислений мы, как правило, делаем много ошибок, при повторении вычислений получаем иногда противоречивые результаты, часто затрудняемся в определении точности, с которой следует вести вычисления и т. п. Вследствие этого нами очень быстро овладевает весьма неприятное чувство неуверенности в правильности окончательных результатов наших вычислений. Для того, чтобы в таких условиях успешно закончить вычисления, необходимо иметь большие знания и навыки, большую опытность во всех вопросах, связанных с практикой вычислительной работы. Эти вопросы подробно излагаются в различных руководствах по приближенным вычислениям, по численному или прикладному анализу и т. п., но и здесь приходится сказать, что действительная опытность в вычислительной работе приобретается также только в результате достаточно большой практики.

Необходимо, далее, иметь в виду, что все чисто арифметические вычислительные операции играют только служебную роль. Но математическая обработка результатов измерений, достаточно полная и всесторонняя, часто позволяет не только оценить правильность и пригодность полученных результатов, но и установить на основании их те физические закономерности, которые имеют место в данном явлении. Для такого всестороннего анализа экспериментальных результатов, который часто является окончательным и очень ответственным завершением обширных лабораторных исследований, были постепенно выработаны особые методы и приёмы; их математической разработкой занимался на протяжении последних двух с половиной столетий ряд крупнейших учёных, главным образом астрономов и математиков, а именно Ньютон, Гаусс, Бессель, Чебышев, Марков, Крылов и многие другие.

Основные вопросы математической обработки экспериментальных результатов в общих чертах становятся понятными из следующих соображений.

1) Все измерения мы можем производить, как известно, лишь с ограниченной точностью; иначе говоря, результат

всякого измерения, с какой бы точностью мы ни выполнили его, даёт нам не точное значение измеряемой величины, а лишь приближённое, т. е. более или менее близкое к её точному значению. Хотя это обстоятельство и не вносит на практике больших осложнений, тем не менее при обработке результатов измерений нам приходится иметь дело с приближёнными значениями различных величин, и обычная школьная арифметика точных чисел оказывается здесь малопригодной. Поэтому при обработке результатов измерений нам приходится пользоваться несколько своеобразной арифметикой приближённых вычислений. Она располагает особыми приёмами, применение которых даёт нам большую экономию времени и несколько не понижает практической точности окончательного результата вычислений. Необходимо напомнить, что с приближёнными вычислениями мы встречаемся очень часто; таковы, например, все вычисления с логарифмами или с тригонометрическими функциями. Точность таких приближённых вычислений определяется числом десятичных знаков в тех таблицах, которыми мы при этом пользуемся. Иначе говоря, точность вычислений здесь можно изменять в зависимости от условий. С приближёнными вычислениями мы встречаемся далее и при других, более сложных вопросах, например при определении корней некоторых уравнений, при вычислении определённых интегралов и т. д. Подобные вычисления мы можем выполнять также с различной степенью точности, т. е. мы можем изменять их точность по нашему желанию.

Приближённые вычисления, с которыми мы встречаемся при обработке результатов измерений, имеют несколько иной характер и часто являются более сложными. Это объясняется главным образом тем, что точность вычислений в этих случаях зависит от точности наших измерений, т. е. она определяется общими условиями измерений, должна быть заранее выяснена и не может быть изменена по нашему желанию. В результате оказывается необходимым более строгий анализ всего вычислительного процесса, а иногда и применение некоторых особых методов и приёмов вычисления. Вследствие этого практическое знакомство со всеми основными приёмами приближённых

вычислений является очень ценным, так как благодаря ему чрезвычайно облегчается и упрощается вся техника вычислений, которая при обычных арифметических приёмах часто могла бы оказаться крайне утомительной, а иногда и непреодолимо сложной.

2) Непосредственно измерять мы можем только очень немного физических величин; такие измерения встречаются лишь в простейших случаях—при измерениях длины, при определении веса тел на обыкновенных весах, при измерениях промежутков времени и т. п. Что же касается всего разнообразия остальных физических величин, то мы их не измеряем непосредственно, а определяем при помощи вычислений, пользуясь различными математическими соотношениями, выведенными на основании законов физических явлений. Прибегать к таким приёмам нам приходится потому, что непосредственные измерения различных физических величин в громадном большинстве случаев могли бы дать недостаточно точные и надёжные результаты или оказались бы очень сложными, а иногда и совершенно невозможными. Имеется очень много случаев, когда мы бываем вынуждены прибегать к таким приёмам косвенных измерений, например:

а) Непосредственные измерения ускорения силы тяжести из наблюдений над свободным падением тел представляются очень сложными и могли бы дать недостаточно точные результаты. Вследствие этого при измерениях g часто применяется так называемый оборотный маятник. Измеряя приведённую длину такого маятника, его период колебаний и при измерениях высокой точности—амплитуду колебаний, вычисляют, пользуясь известной формулой маятника, ускорение силы тяжести.

б) Электрическое сопротивление различных тел очень часто определяется при помощи мостика сопротивлений. Здесь нам приходится непосредственно измерять длину двух частей измерительной проволоки мостика, а затем вычислять искомое сопротивление по формуле, которую мы выводим теоретически на основании законов электродинамики.

в) Длину световой волны часто определяют при помощи дифракционной решётки. При этом нам прихо-

дится измерять углы отклонения спектральных линий от начального направления лучей, а затем вычислять соответствующие им длины волн, пользуясь формулой дифракционной решётки.

Во всех подобных случаях мы подставляем в те или иные формулы результаты наших измерений, т. е. величины приближённые. Отсюда следует, что приближённым должен быть и результат наших вычислений, т. е. значение искомой величины. При этом ошибка результата, очевидно, должна определяться не только точностью наших измерений, но и видом данной формулы. Последнее обстоятельство поставило перед нами вопрос о зависимости ошибок наших вычислений от вида формул, которые служат для вычислений. При общей математической разработке этих вопросов применяют приёмы дифференциального исчисления, считая, что искомая величина является функцией величин, непосредственно измеряемых, причём вид функциональной зависимости определяется математическим выражением соответствующего физического закона. Таким образом, с математической стороны вопрос сводится к тому, чтобы установить приёмы, которые дают возможность вычислять ошибки функций, если даны ошибки их аргументов и вид функциональной зависимости.

Все эти вопросы были решены вполне удовлетворительно, т. е. мы можем находить ошибки любых функций, зная ошибки их аргументов. Одновременно оказалось возможным решить и ещё один вопрос, который хотя и не имеет прямого отношения к обработке результатов измерений, но всё же тесно с ней связан. Это—вопрос об определении наилучших условий для измерений, т. е. таких условий, при которых ошибка результата должна быть наименьшей.

3) Каждое физическое измерение, как известно, всегда повторяют несколько раз, затем вычисляют среднее арифметическое из всех результатов, полученных при отдельных измерениях, и считают, что таким приёмом мы находим значение измеряемой величины. Этот обычный приём определения окончательного результата измерений является, несомненно, очень простым, и при его применении вряд ли можно ожидать каких-либо затруднений.

Только одно небольшое осложнение может возникнуть здесь вследствие необходимости установить ту точность, с которой следует производить вычисление среднего арифметического, но и это затруднение легко устраняется. Однако, несмотря на всю простоту нахождения среднего арифметического, всё же оказалось, что для окончательной оценки результатов измерений необходимо пользоваться выводами теории случайных ошибок, которая была разработана главным образом Гауссом. Эта теория позволила математически обосновать постулат среднего арифметического, но вместе с тем привела к выводу, что среднее арифметическое является лишь некоторым приближением к точному значению измеряемой величины. Хотя теория даёт возможность вычислить степень этого приближения, однако точное значение измеряемой величины всё же продолжает оставаться неизвестным. Этот на первый взгляд непонятный результат является следствием одной особенности теории случайных ошибок: положения и выводы, основанные на этой теории, являются лишь вероятными, что объясняется особыми свойствами тех ошибок, которые мы неизбежно допускаем при измерениях. Вследствие этого и окончательный результат измерений, вычисленный по этой теории, приводит лишь к наиболее вероятному значению измеряемой величины, точное значение которой, как уже было сказано, всё же остаётся нам неизвестным.

Не следует, однако, думать, что теория случайных ошибок вследствие этой особенности является мало удовлетворительной или что она не даёт надёжного метода для нахождения окончательных результатов измерений. Наоборот, эта теория представляет собой очень совершенное построение, все её выводы являются практически чрезвычайно ценными и находят самое широкое применение как в физике, так и в других науках, где приходится иметь дело с измерениями или наблюдениями, например, в астрономии, в механике, в химии, в различных отраслях научной и практической техники и т. д. Не менее широкие области применения находит эта теория и в таком большом отделе прикладной математики, как математическая статистика.

4) Результаты измерений очень часто изображают графически, причём в большинстве случаев пользуются системой обычных прямоугольных координат. Однако этот способ графического представления функциональных зависимостей решает далеко не все вопросы, возникающие при обработке экспериментальных результатов. Поэтому постепенно были разработаны другие, очень разнообразные и иногда более сложные приёмы графических построений, в основу которых положено весьма полное и всестороннее исследование геометрических и аналитических свойств самых разнообразных функций. Все эти разнообразные методы искусственных графических построений, разработкой которых занимается особая отрасль прикладной математики, так называемая номография, находят в настоящее время очень широкое применение и в физике и в других науках, например, в технических.

5) При изучении различных физических явлений очень часто приходится исследовать зависимость одной какой-либо величины от некоторых других величин, например, от времени, температуры, давления, концентрации и т. п., иначе говоря, приходится систематически определять значения одной величины при различных значениях другой величины. В результате получается ряд соответственных значений той и другой величины, что, вообще говоря, даёт возможность выразить зависимость между этими величинами в виде некоторого математического соотношения. Но вместе с тем такие ряды соответственных значений двух величин очень часто показывают настолько сложную зависимость между величинами, что установить её точное математическое выражение, т. е. подобрать соответствующую функцию, не представляется возможным. В таких случаях вместо точной функции, которая нам остаётся неизвестной, подбирают другую функцию, приближённую, более простого вида, которая даёт возможность определять приближённые значения одной из величин при любых значениях другой величины. Эти вопросы решаются методами интерполяции или интерполяции. При более полном систематическом изучении различных физических явлений с этими приёмами приходится очень часто встречаться.

Несколько своеобразные приёмы интерполирования были разработаны для тех случаев, когда зависимость между исследуемыми величинами носит периодический характер. В этих случаях также вводят приближённую периодическую функцию, подобранную так, чтобы её можно было разложить на некоторое число простых синусоидальных колебаний. Методы разложения периодических функций на синусоидальные составляющие рассматриваются в особом разделе прикладной математики, так называемом гармоническом анализе, который находит очень широкое применение в некоторых разделах физики, а также в таких науках, как учение о переменных токах, радиотехника и т. д.

6) С интерполированием тесно связаны вопросы о составлении так называемых эмпирических функций или эмпирических формул. В эти формулы кроме исследуемых величин входят ещё один или, чаще, несколько коэффициентов, числовые значения которых устанавливаются на основании того же экспериментального материала. Эмпирические формулы не вполне отвечают действительным законам физических явлений, так как представляют собой формулы приближённые; кроме того, эмпирические формулы, как правило, можно применять только в ограниченном интервале, обычно в пределах данной серии измерений. Тем не менее эмпирические формулы пользуются очень широким распространением, так как они дают прекрасное обобщение экспериментальных материалов и часто являются первой, весьма важной ступенью на пути к установлению более глубоких закономерностей, а иногда и общих физических законов.

7) При обработке результатов измерений очень длительным и трудным процессом является собственно вычислительная работа, т. е. чисто арифметические операции. Для облегчения этой работы применяются различные математические таблицы, большинство которых общеизвестно. С той же целью ускорения и улучшения вычислений уже давно были разработаны приёмы производства вычислений механическим путём при помощи соответствующих приборов. К числу таких приборов принадлежат различные счётные линейки, например, обыкновенная

логарифмическая линейка, и различные счётные машины, например арифмометры; в последнее время начинают применяться особые вычислительные приборы, основанные на свойствах электронно-лучевых трубок. Точно так же при обработке результатов измерений часто приходится прибегать к графическим вычислениям, таким, например, как измерение площадей, ограниченных некоторыми кривыми, аналитическое выражение которых может оставаться неизвестным. Такие операции в настоящее время выполняются, как правило, также при помощи различных механических и электронных приборов, которые носят название планиметров, интеграторов, интеграфов. К числу таких же приборов принадлежат и так называемые гармонические анализаторы, т. е. интеграторы специального типа для разложения периодических кривых на гармонические составляющие. Со всеми подобными приборами, действие которых основано обыкновенно на очень остроумных и интересных математических соображениях, следует практически познакомиться.

Тем, что было изложено выше, далеко не ограничивается круг вопросов, которые можно встретить при обработке экспериментальных результатов, и очень часто экспериментатору приходится разрешать ряд других вычислительных задач. Так, очень часто приходится решать систему уравнений приближённо, вычисляя их корни с определённой степенью точности; также часто приходится приближённо вычислять определённые интегралы, если их значение не удаётся найти обычными приёмами интегрального исчисления, и т. д. Для таких вычислительных операций разработаны определённые приёмы, иногда весьма интересные. Но эти вопросы всё же несколько выходят за пределы того, что составляет предмет математической обработки результатов измерений, по крайней мере в её элементарной форме, и вследствие этого в дальнейшем не рассматриваются; их изложение можно найти в более подробных курсах по вычислительному или прикладному математическому анализу.

ГЛАВА I ПРИБЛИЖЁННЫЕ ЗНАЧЕНИЯ ВЕЛИЧИН И ИХ ОШИБКИ

1. Приближённые значения величин. Мы встречаем очень много таких чисел и величин, точное значение которых остаётся нам неизвестным; в таких случаях вместо точных значений этих величин мы пользуемся их значениями, не вполне точными, или, как принято говорить, их **п р и б л и ж ё н н ы м и з н а ч е н и я м и**. Иногда числа и величины такого рода называют для краткости также **п р и б л и ж ё н н ы м и ч и с л а м и** и **п р и б л и ж ё н н ы м и в е л и ч и н а м и**, хотя это и является не вполне правильным, так как следует точнее говорить только о приближённых значениях таких величин.

Приближённые значения мы вводим при вычислениях чрезвычайно часто. Так, приближёнными значениями мы всегда пользуемся для таких чисел, как π , или основание натуральных логарифмов e , приближённые значения мы применяем также для громадного большинства тригонометрических функций, логарифмов и т. п. Приближённые значения величин мы получаем и в результате всех физических измерений, как об этом уже было сказано раньше (стр. 11). Отсюда видно, что в физике мы встречаемся с приближёнными значениями величин очень часто, значительно чаще, чем с числами точными, в чём очень нетрудно убедиться. Рассмотрим несколько примеров.

Пример 1. Объём V цилиндра мы можем определить, если измерим его высоту h и радиус основания r , а затем воспользуемся обычной формулой:

$$V = \pi r^2 h.$$

В этой формуле, как мы видим, только показатель степени является числом точным, а для всех остальных величин мы имеем приближённые значения. Действительно, для π мы берём одно из его приближённых значений, r и h находим из измерений, т. е. для них получаются тоже приближённые значения. Отсюда следует, что и объём цилиндра V мы можем определить тоже только приближённо.

Пример 2. Частоту ν линий водородного спектра можно вычислить, пользуясь формулой

$$\nu = \frac{2\pi^2 m e^4}{h^3} \left(\frac{1}{k^2} - \frac{1}{n^2} \right),$$

где m и e обозначают массу и заряд электрона, h обозначает постоянную Планка, а k и n обозначают номера устойчивых орбит электрона в водородном атоме. Хотя в этой формуле имеется много точных чисел, а именно 2, k , n и все показатели степеней, однако для величин π , m , e и h мы имеем приближённые значения, и в результате вычислений для величины ν получается также приближённое значение.

Пример 3. При вычислении коэффициентов поглощения k квантового излучения различных видов (лучи света, рентгеновские лучи, γ -лучи) применяется формула

$$k = \frac{\ln I_0 - \ln I}{d},$$

где I_0 и I обозначают интенсивность излучения до и после прохождения его через поглощающий слой, толщина которого равна d . В этой формуле для всех величин мы имеем приближённые значения.

Аналогичные условия мы встречаем и во всех других многочисленных и разнообразных физических формулах: приближённых чисел в них оказывается значительно больше, чем точных.

2. Округление приближённых значений чисел; правило дополнения. Приближённые значения чисел и величин мы можем брать с различным числом десятичных знаков, т. е. с различной точностью, или

с р а з л и ч н ы м п р и б л и ж е н и е м. Так, число π часто принимают равным 3,14159, нормальное значение ускорения силы тяжести g_0 принято считать равным $980,665 \text{ см сек}^{-2}$, массу атома водорода принимают равной $1,673 \cdot 10^{-24} \text{ г}$ и т. д. Если такая точность оказывается излишне большой, то те же величины можно брать с меньшим числом десятичных знаков—округлять приближённые числа, отбрасывая излишние знаки. При этой операции пользуются простым правилом дополнения, т. е. увеличивают на единицу последнюю из остающихся цифр, если первая из отбрасываемых цифр больше 5. Так, для приведённых выше величин можно, ограничиваясь меньшей точностью, принять, например: 3,1416 (отбрасывается цифра 9), или 3,14 (отбрасываются цифры 159), 981 см сек^{-2} (отбрасываются цифры 665) и $1,7 \cdot 10^{-24} \text{ г}$ (отбрасываются цифры 73).

Для тех случаев, когда первая из отбрасываемых цифр равна 5, общепринятого правила дополнения не установлено, и на практике пользуются двумя различными приёмами.

1) Первую из остающихся цифр увеличивают на единицу, за исключением тех случаев, когда отбрасываемая цифра 5 сама появилась в результате применения правила дополнения, что иногда может иметь место; в таких случаях, отбрасывая цифру 5, последнюю из остающихся цифр оставляют без изменения. Так, например, из математических таблиц мы находим, что длина дуги окружности, отвечающей углу 65° , с точностью до шестого десятичного знака равна 1,134464 длины радиуса окружности. Округляя это число до четвёртого десятичного знака, следует взять 1,1345; если в этом числе мы решим отбросить ещё и четвёртый десятичный знак, то следует взять 1,134, а не 1,135. Очевидно, что этот приём является достаточно обоснованным.

2) Первую из остающихся цифр увеличивают на единицу, если эта цифра нечётная, и оставляют её без изменения, если она чётная. Так, округляя π до третьего десятичного знака, следует взять 3,142, но, округляя числа 12,665 или 10,385 до второго десятичного знака, следует взять 12,66 и 10,38.

Первое из этих правил находит, повидимому, более широкое применение. Во всяком случае следует выбрать определённый приём округления чисел и точно следовать ему при всех приближённых вычислениях; это значительно снижает те ошибки, которые мы неизбежно вносим при округлении чисел.

Если при округлении некоторого числа последнюю остающуюся цифру не приходится увеличивать, то новое значение числа, очевидно, меньше его прежнего значения; в этих случаях говорят, что число взято с некоторым недостатком. Если же при округлении приближённого числа последнюю остающуюся цифру приходится увеличить и новое значение числа оказывается больше его прежнего значения, то говорят, что приближённое число взято с некоторым избытком. Величину недостатка или избытка, которые появляются в результате округления чисел, нетрудно, очевидно, приближённо вычислить. Так, в приведённых выше примерах значение числа π , принятое равным 3,14, взято с недостатком, приближённо равным 0,00159, а значения π , ускорения силы тяжести и массы водородного атома, принятые равными соответственно 3,142, 981 см сек^{-2} и $1,7 \cdot 10^{-24} \text{ г}$, все взяты с избытком, приближённо равным: для первой величины 0,00041, для второй 0,335 см сек^{-2} и для третьей $0,027 \cdot 10^{-24} \text{ г}$.

3. Ошибки приближённых значений. Приближённые значения всегда имеют некоторую погрешность, или ошибку величину которой принято определять, во-первых, абсолютной погрешностью, или абсолютной ошибкой, приближённых значений и, во-вторых, их относительной погрешностью, или относительной ошибкой.

1) Абсолютной ошибкой приближённого значения некоторой величины называют разность между точным и приближённым значениями этой величины. Таким образом, можно написать:

$$\varepsilon = X - A,$$

где ϵ обозначает абсолютную ошибку приближённого значения A некоторой величины, точное значение которой равно X .

Из этой формулы находим:

$$X = A + \epsilon, \quad (1)$$

т. е. точное значение некоторой величины X получается в результате сложения приближённого значения этой величины и его абсолютной ошибки.

2) Относительной ошибкой приближённого значения некоторой величины называют отношение его абсолютной ошибки к точному значению данной величины. Таким образом, можно написать:

$$\delta = \frac{\epsilon}{X}, \quad (2)$$

где ϵ и X имеют прежние значения, а δ обозначает относительную ошибку приближённого значения A величины X .

Величина X в формулах (1) и (2), т. е. точное значение этой величины, остаётся нам неизвестной. Отсюда следует, что определить точные значения величин ϵ и δ , т. е. абсолютной и относительной ошибок, мы тоже не можем. Это было бы возможно в том случае, если бы величина X в уравнениях (1) и (2) была нам точно известна. Такие случаи иногда могут иметь место, например, если мы, желая несколько упростить вычисления, округляем значения каких-либо точных чисел, т. е. берём для них более простые приближённые значения. Например: вместо дроби 0,3125 можно при вычислениях небольшой точности взять приближённо 0,31 или даже 0,3. В таких случаях по формулам (1) и (2) можно точно вычислить значения абсолютной и относительной ошибок, которые мы при этом допускаем. Все подобные случаи, очень простые и относительно редкие, в дальнейшем мы исключаем из рассмотрения, т. е. величину X мы будем считать приближённой; иначе говоря, точные значения величин X , ϵ и δ в уравнениях (1) и (2) будем считать неизвестными. Мы останавливаемся на таком предположении потому, что по отношению к результатам физических измерений оно всегда имеет место.

Далее мы будем предполагать, что абсолютная ошибка ϵ очень мала по сравнению с величиной X и с её приближённым значением A . Это предположение по отношению к физическим измерениям также оправдывается, но только в тех случаях, когда измерения выполняются с достаточно большой точностью, т. е. по отношению к так называемым точным физическим измерениям, которые и предполагаются в дальнейшем.

При этом условии, во-первых, имеют место неравенства

$$\left. \begin{aligned} \epsilon &\ll X \\ \epsilon &\ll A \end{aligned} \right\} \quad (3)$$

и, во-вторых, величины X и A следует считать очень близкими одна к другой, так как их разность, равная ϵ , является очень малой по сравнению с величинами X и A .

Вследствие этого в выражении для относительной ошибки (2) величину X можно заменить близкой к ней величиной A , т. е. положить, что относительная ошибка δ приближённо определяется выражением

$$\delta = \frac{\epsilon}{A}. \quad (4)$$

На основании этой формулы, преобразуя правую часть выражения (1), находим:

$$X = A \left(1 + \frac{\epsilon}{A} \right) = A(1 + \delta),$$

где δ обозначает приближённое значение относительной ошибки, причём очевидно, что при условиях (3) имеет место неравенство

$$\delta \ll 1.$$

Отсюда мы получаем такой вывод:

Если абсолютная ошибка приближённого значения некоторого числа мала по сравнению с его величиной, то относительная ошибка этого приближённого значения мала по сравнению с единицей.

Относительные ошибки приближённых значений принято выражать в процентах; вследствие этого относительные

ошибки иногда называют процентными ошибками. Таким образом, обозначая относительную ошибку, выраженную в процентах, или процентную ошибку, через $\delta\%$, мы можем написать:

$$\delta\% = 100\delta = 100 \frac{\varepsilon}{A}.$$

В дальнейшем будем считать, что относительная ошибка в общем случае всегда выражается в процентах. Но при вычислениях обыкновенно приходится переходить от относительной ошибки, выраженной в процентах, к её значению, определяемому формулой (4). Этот переход выполняется простым делением на 100 относительной ошибки, выраженной в процентах.

Из формулы (4) находим, что

$$\varepsilon = A\delta. \quad (5)$$

Таким образом, мы видим, что абсолютная ошибка приближённого значения некоторой величины равна произведению этого приближённого значения на его относительную ошибку.

Для определения относительных ошибок δ приближённых значений нам необходимо, очевидно, знать и их абсолютные ошибки. Действительно, при переходе от формулы (2) к формуле (4) мы заменили неизвестную нам величину X её приближённым значением A , но числитель ε в формуле (4), т. е. точное значение абсолютных ошибок, продолжает оставаться нам неизвестным. Вследствие этого нам приходится вместо точных значений абсолютных ошибок ограничиваться их приближёнными значениями, а затем по формуле (4) находить соответствующие приближённые значения относительных ошибок. Необходимо при этом различать два случая.

1) Для таких величин, значения которых известны с большим числом десятичных знаков, например, для таких чисел, как π , e , логарифмы и т. п., приближённую величину абсолютных ошибок всегда можно непосредственно вычислить с достаточной точностью. Так, например, значение числа π мы можем взять из таблиц с доста-

точным числом знаков

$$\pi = 3,14159265\dots$$

Отсюда находим непосредственными вычислениями для различных приближений π :

π	ε	δ
3,1	+0,04159265	+1,3%
3,14	+0,00159265	+0,05%
3,142	—0,00040735	—0,01%
3,1416	—0,00000735	—0,0002%
3,14159	+0,00000265	+0,00008%
3,141593	—0,000000035	—0,00001%

Второй столбец даёт с точностью до восьмого десятичного знака приближённые значения абсолютных ошибок различных приближений π . В третьем столбце приведены приближённые значения относительных ошибок также для различных приближений π , вычисленные по формуле (4). То же самое можно, очевидно, применить ко всем приближённым значениям таких величин, которые известны с достаточно большой точностью.

2) Но если мы имеем результаты измерений, то непосредственно вычислить их абсолютные ошибки не представляется возможным. Вообще говоря, мы можем только утверждать, что при правильных и тщательно выполненных измерениях абсолютные ошибки их результата не должны выходить за пределы так называемой точности измерений, т. е. той наименьшей величины, которую мы можем определять при измерениях с уверенностью в правильности получаемых результатов; иными словами, мы можем только утверждать, что значение абсолютной ошибки окончательного результата каких-либо измерений приближённо должно определяться точностью этих измерений. Так, например, точность измерений длины при помощи винтового микрометра обыкновенно бывает равна 0,01 мм. Если для таких же измерений применять так называемый компаратор, то точность измерений нетрудно повысить до 1 м. Наконец, интерференционные методы измерения длины могут дать точность до 0,1 м. Можно утверждать, что абсолютные ошибки результатов

измерений в этих трёх случаях не должны выходить за пределы соответственно $\pm 0,01$ мм, ± 1 м и $\pm 0,1$ м; иначе говоря, эти величины при правильно произведённых измерениях определяют собой приближённые значения абсолютных ошибок результата соответственно в первом, втором и третьем случаях.

Хотя впоследствии мы познакомимся с другим, более точным методом определения абсолютных ошибок измерений (гл. V), однако из того, что было только что сказано, следует, что точность измерений всегда необходимо предварительно определять, т. е. что, приступая к каким-либо измерениям, необходимо прежде всего определить ту точность, которая может быть достигнута в данном случае, так как при этом условии можно всегда приближённо установить возможное значение абсолютной ошибки результата измерений.

4. Знак ошибок приближённых значений. Необходимо, далее, обратить внимание на те осложнения, которые возникают при определении знака ошибок приближённых значений. Вновь приходится указать на то, что для приближённых значений тех величин, которые известны с большим числом десятичных знаков, т. е. для таких величин, как π , e и т. п., мы можем непосредственно определять не только приближённую величину их абсолютных и относительных ошибок, но и знак этих ошибок, как это видно из приведённой выше таблицы приближённых значений π . Таким образом, в этих случаях знак ϵ и δ в формулах (1) и (4) нам всегда известен. Что же касается результатов физических измерений, то, хотя возможно, как мы видим, установить приближённое значение их абсолютных ошибок, однако знак ошибок остаётся нам неизвестным. Это следует из того, что при измерениях мы никогда не можем определить, допускаем ли мы ошибку в сторону некоторого преуменьшения измеряемой величины (знак ϵ в формуле (1) положителен) или в сторону её преувеличения (знак ϵ отрицателен). Вследствие этого, применяя формулы (1), (4) и (5) к результатам измерений, необходимо иметь в виду, что ϵ и δ могут иметь и положительное и отрицательное значения; таким образом, для

этого случая следует написать:

$$X = A \pm \epsilon, \quad (1')$$

$$\delta = \pm \frac{\epsilon}{A}, \quad (4')$$

$$\epsilon = \pm A\delta. \quad (5')$$

Это обстоятельство в теории ошибок играет очень большую роль, и при всех измерениях необходимо иметь в виду, что ϵ и δ могут быть и положительными и отрицательными, причём их знак остаётся нам неизвестным.

Следует указать ещё на некоторые черты различия между абсолютной и относительной ошибками, проявляющиеся в основном в двух отношениях.

1) Так как в результате измерений мы получаем значения различных физических величин, которые в громадном большинстве случаев имеют определённую размерность, то и абсолютные ошибки, как видно из формулы (1'), являются величинами размерными, и их размерность отвечает размерности измеряемой величины. Что же касается относительных ошибок, то, как видно из формулы (4'), они определяются отношением двух величин одной и той же размерности; вследствие этого относительные ошибки всегда являются величинами безразмерными. Благодаря этому относительные ошибки дают возможность сравнивать точность физических измерений, совершенно разнородных по существу, например, таких, как измерение масс, измерение длин, измерение времени и т. д. Необходимость в таком сравнении точности разнородных измерений очень часто встречается на практике (гл. III).

2) Достоинство или качество измерений зависит от их точности, т. е. от величины ошибок, допущенных при этих измерениях: чем больше точность измерений, тем меньше их ошибки, тем выше и качество измерений. Но при оценке качества измерений абсолютная и относительная ошибки оказываются неравноправными, и нетрудно убедиться в том, что особенно наглядное представление о качестве измерений нам даёт их относительная ошибка. Это становится понятным из следующих соображений.

Обе ошибки, абсолютная ϵ и относительная δ , как сказано, зависят от точности измерений, но относительная

ошибка δ , как видно из формулы (4'), зависит, кроме того, от размеров измеряемой величины A . Отсюда следует, что при одной и той же абсолютной ошибке мы можем получить совершенно различные по величине относительные ошибки, если измеряемые величины сильно отличаются по своим размерам, так как чем больше измеряемая величина, тем меньше становится относительная ошибка при той же величине абсолютной ошибки. Но чем больше измеряемая величина, тем сложнее, вообще говоря, производить измерения, сохраняя прежнюю величину абсолютной ошибки, и если этого всё же удалось достичь, то измерения следует считать высокими по качеству. Одновременно относительная ошибка таких измерений становится очень малой, хотя их абсолютная ошибка, как сказано, сохраняет свою прежнюю величину. Отсюда и следует, что более наглядным показателем качества измерений служит их относительная ошибка.

Рассмотрим несколько примеров.

Пример 1. Если мы измеряем толщину пластинок при помощи винтового микрометра с точностью до $0,01$ мм, то при всех измерениях абсолютная ошибка результата остаётся одной и той же, приближённо равной $\pm 0,01$ мм. Но относительные ошибки измерения двух пластинок с толщиной, приближённо равной 1 мм и 1 см, оказываются совершенно различными: относительная ошибка измерения первой пластинки достигает $\pm 1\%$, тогда как для второй пластинки она составляет только $0,1\%$, т. е. оказывается в десять раз меньше. Таким образом, мы видим, что толщина первой пластинки измерена с большей относительной ошибкой и что измерение толщины второй пластинки с качественной стороны является более высоким. Отсюда следует также, что если нам необходимо измерить толщину первой пластинки с той же относительной ошибкой, как и толщину второй пластинки, то для этого, очевидно, необходимо применить более точный прибор, например сферометр.

Пример 2. Масса атома водорода принимается равной $(1,673 \pm 0,001) \cdot 10^{-24}$ г, масса электрона принимается равной

$(9,11 \pm 0,01) \cdot 10^{-28}$ г. Отсюда следует, что абсолютная ошибка определения массы электрона, равная $\pm 10^{-30}$ г, в тысячу раз меньше абсолютной ошибки определения массы атома водорода, т. е. $\pm 10^{-27}$ г, и могло бы казаться, что масса электрона определена с исключительной точностью. Но относительные ошибки того и другого измерений, равные соответственно $\pm 0,06\%$ и $\pm 0,11\%$, показывают, что масса водородного атома нам известна с относительно большей точностью.

Пример 3. Последние измерения скорости света в пустоте дали число $(2,99796 \pm 0,00004) \cdot 10^{10}$ см/сек, точные измерения скорости звука в воздухе при 0°C привели к числу $(331,63 \pm 0,04)$ м/сек. Абсолютные ошибки этих величин, выраженные в см/сек, приближённо равны: для скорости света $\pm 4 \cdot 10^5$ см/сек и для скорости звука ± 4 см/сек; иначе говоря, абсолютная ошибка скорости света в пустоте в 100 000 раз больше абсолютной ошибки скорости звука в воздухе. Но относительная ошибка первой величины оказывается равной $\pm 0,001\%$, тогда как для второй величины она достигает $0,012\%$.

Таким образом, мы вправе утверждать, что скорость света в пустоте измерена со значительно большей точностью.

Эти примеры и изложенные выше соображения показывают, что вычисление относительных ошибок очень часто может оказаться весьма полезным. Вследствие этого принято, вычисляя окончательный результат каких-либо измерений, указывать обе его ошибки, абсолютную и относительную.

5. Классификация приближённых чисел. Выше уже были указаны некоторые черты различия между приближёнными значениями таких величин, как π , e , ..., с одной стороны, и результатами измерений, с другой. Рассматривая этот вопрос более подробно, нетрудно прийти к выводу, что приближённые значения всех величин по отношению к их ошибкам и влиянию последних на результат вычислений можно разделить на три типа:

1) В о-п е р-в-ы-х, приближённые значения таких величин, как π , e , $\sqrt{2}$, логарифмы и т. п. Все эти величины

нам известны с очень большой точностью; их значения даются в различных математических таблицах.

2) В о-в т о р ы х, приближённые значения различных физических постоянных, например, значения плотности тел, их коэффициентов расширения, значения удельного сопротивления, длин волн спектральных линий и т. п. Все эти величины нам известны в большинстве случаев также с очень большой точностью; их значения даются в различных таблицах физических величин.

3) В-т р е т ь и х, результаты обычных лабораторных измерений, при которых мы также находим приближённые значения различных физических величин, но в громадном большинстве случаев со значительно меньшей точностью.

По отношению к приближённым числам этих трёх типов можно сделать следующие замечания.

1) Приближённые значения первого типа мы всегда берём с некоторым округлением; иначе говоря, мы можем их брать с различной точностью, практически, сколь угодно большой; соответственно этому изменяется величина их ошибок, причём знак и величину ошибок мы всегда можем определить, как это было показано на примере числа π (стр. 24). Возможность изменять по нашему усмотрению величину ошибок приближённых значений этого типа позволяет считать их практически не приближёнными, а точными числами. Действительно, при любых вычислениях мы можем вводить значения таких величин с тем приближением, которое в данном случае необходимо; иначе говоря, мы всегда можем сделать достаточно малыми те ошибки, которые могли бы внести эти приближённые значения в результаты вычислений. Отсюда, однако, не следует, что все приближённые значения таких величин, как π , e , $\sqrt{2}$ и т. п., надо вводить при вычислениях всегда с большим числом десятичных знаков: последнее в каждом отдельном случае может быть различным, так как оно определяется точностью наших вычислений, т. е. величиной той ошибки, которую мы допускаем в окончательном результате вычислений; эту ошибку всегда необходимо предварительно установить. Все эти вопросы решаются при

помощи ряда приёмов, которые мы рассмотрим несколько позднее.

2) Приближённые значения второго типа, т. е. принятые в настоящее время значения различных физических постоянных, были установлены в большинстве случаев в результате весьма точных измерений, часто в результате очень большого числа специальных измерительных работ, чрезвычайно тщательно выполненных. Вследствие этого значения данных величин нам известны с очень большой точностью, и их ошибки в громадном большинстве случаев значительно меньше тех, которые мы по необходимости допускаем при обычных измерительных работах в лабораториях. Таким образом приближённые значения этого типа очень часто можно считать практически точными числами, подобно тому как это имеет место по отношению к приближённым значениям первого типа, от которых они отличаются только своей несколько меньшей точностью, а также тем, что знак их ошибки, вообще говоря, остаётся нам неизвестным. Только в тех случаях, когда мы производим измерения и вычисления с очень большой точностью, необходимо считаться с ограниченной точностью приближённых значений различных физических величин и с двойным знаком их ошибок.

3) Приближённые значения третьего типа, т. е. результаты обычных измерительных работ в лабораториях, являются главным предметом всего дальнейшего анализа. Рассматривая результаты измерений как приближённые числа, необходимо сделать следующее весьма важное разъяснение, которое, конечно, имеет отношение и к приближённым значениям второго типа.

Причины, в силу которых при всех измерениях мы неизбежно допускаем некоторые ошибки (стр. 11), могут быть очень различными: неправильные или неточные показания приборов, влияние внешних факторов, неправильно произведённые отсчёт или запись и т. п. Кроме того, при измерениях неизбежно оказываются и такие ошибки, причины которых остаются неясными, неопределёнными, нам неизвестными. Исследование всех этих вопросов позволило разделить все ошибки измерений на два вида: во-первых, о ш и б к и с и с т е м а т и ч е с к и е и,

во-вторых, ошибки случайные. Иногда указывают на третий вид ошибок: неправильная запись наблюдений, описки, просчёты и т. д., но таких грубых ошибок мы в дальнейшем касаться не будем.

а) Систематические ошибки вызываются главным образом неправильными показаниями приборов или ошибочностью метода измерений, или постоянным, но односторонним внешним воздействием. Так, измерение температуры термометром со смещённой нулевой точкой будет систематически неправильным, пока в результаты измерений мы не внесём соответствующей поправки; точно так же неравномерное нагревание коромысла весов, например лучами Солнца или теплотой радиатора, будет давать систематическую ошибку при взвешивании. Систематические ошибки, как правило, оказывают при измерениях постоянное влияние, и, несмотря на это, обнаруживать и исключать их является очень сложной задачей. Это достигается в результате чрезвычайно тщательного изучения всего метода измерений и отдельных приборов, их проверкой и исправлением и, если оказывается необходимым, введением соответствующих поправок в результаты измерений; очень часто приходится также сопоставлять получаемые результаты с теми, которые были получены при аналогичных работах другими исследователями.

б) Случайные ошибки измерений, обычно очень небольшой величины, являются следствием тех незначительных неточностей, которые мы неизбежно делаем при установке приборов и отсчётах их показаний. Случайные ошибки всегда наблюдаются при измерениях. Действительно, практика показывает, что если мы повторяем какое-либо измерение несколько раз, то, как бы тщательно мы ни производили все операции, всё же получаются несколько различные результаты. Величина случайных ошибок отдельных измерений не должна во много раз превышать точности измерений (стр. 25), например, величина случайных ошибок отдельных измерений может превышать их точность в два, три, четыре раза, но не в десять или более раз. В последнем случае правильность выполнения измерений становится сомнительной.

Хотя исключить случайные ошибки при измерениях не представляется возможным, однако математическая теория ошибок указывает приёмы, которые позволяют уменьшить влияние этих ошибок на окончательный результат измерений и более точно установить его ошибку (гл. IV и V). Применять математическую теорию ошибок можно исключительно к случайным ошибкам, для которых она и была разработана. Вследствие этого во всём дальнейшем материале мы всегда будем предполагать, что все систематические ошибки, так же как и грубые просчёты, совершенно исключены при измерениях, и мы имеем дело исключительно с случайными ошибками.

6. Верные цифры в приближённых значениях чисел. Значащими цифрами в числах принято называть все цифры 1, 2, 3, ..., 9, а также нуль, но только в тех случаях, если он стоит в середине или на конце числа; если же нули стоят в десятичной дроби с левой стороны для указания разряда остальных цифр, то они значащими цифрами не считаются. Так, если мы возьмём такие числа:

$$\pi = 3,14159,$$

$$e = 2,71828,$$

$$g_0 = 980,665 \text{ см сек}^{-2},$$

$$1 \text{ сутки} = 86\,400 \text{ сек.},$$

$$1 \text{ абс.эл.-ст.ед.пот.} = 300 \text{ в},$$

то первые три числа имеют по шести значащих цифр, четвёртое—пять, а последнее—три значащие цифры.

Одновременно, если мы напишем дроби: 1,017; 0,17; 0,017 и 0,0017, то первая из них имеет четыре значащие цифры, а остальные—по две значащие цифры. Последнее становится ясным, если эти три дроби написать в таком виде: $17 \cdot 10^{-2}$, $17 \cdot 10^{-3}$ и $17 \cdot 10^{-4}$.

В арифметике приближённых вычислений кроме понятия о значащих цифрах приходится вводить ещё понятие о верных цифрах, или верных знаках

приближённого числа. Необходимо с самого начала указать на то, что в приближённых числах все цифры должны быть верными, за исключением только одной последней цифры или последнего знака, относительно которого можно сделать следующие замечания.

1) В приближённых числах первого типа (стр. 29) последний знак является верным, как это можно непосредственно видеть из таблицы значений π . Действительно, ни одну из последних цифр во всех значениях π в этой таблице мы не можем заменить никакой другой цифрой, так как ошибки тех значений π , которые могли бы получиться при такой замене, должны были бы резко возрасти. То же самое имеет место по отношению ко всем числам, которые мы берём с округлением, например по отношению к приближённым числам второго типа, взятым с округлением. Таким образом, можно сказать, что если приближённые числа берутся с округлением, то все их знаки, не исключая и последнего, являются верными. Отсюда также следует, что абсолютная ошибка таких приближённых чисел не может быть больше половины, или $\pm 0,5$ от одной единицы последнего знака приближённого числа. Справедливость этого вывода подтверждается той же таблицей значений π .

2) В приближённых числах третьего типа, а также в числах второго типа, если мы их берём без округления, последний знак является сомнительным. Так как возможные изменения последнего знака приближённых чисел определяются их абсолютными ошибками, то последние должны быть известны. Принимая во внимание величину абсолютных ошибок приближённых чисел, мы можем:

во-первых, установить общие приёмы для правильной записи окончательных результатов физических измерений и,

во-вторых, найти простые приёмы: 1) для быстрого определения относительных ошибок приближённых чисел в зависимости от числа их верных знаков, 2) для решения обратной задачи, т. е. для определения числа верных знаков в приближённых числах в зависимости от величины их относительных ошибок.

1) Запись приближённых значений чисел. Результаты измерений очень часто дают в таком виде, что их абсолютная ошибка оказывается равной ± 1 от последнего десятичного знака, например очень много физических констант в таблицах даётся с такими пределами их абсолютных ошибок. При этом условии величину абсолютной ошибки очень часто не указывают. Однако на практике мы часто встречаемся с такими случаями, когда следовать этому правилу не представляется возможным, так как абсолютные ошибки окончательного результата измерений часто выходят, иногда очень значительно, за указанные пределы. Всё это вызывает необходимость при записи окончательных результатов измерений следовать некоторым правилам, которые проще всего можно разъяснить на ряде примеров.

Пример 1. Основной единицей для измерения плотностей тел принята, как известно, плотность дистиллированной воды при 4°C . Весьма точные измерения и вычисления плотности воды при других температурах дали возможность установить эти величины с точностью до шестого десятичного знака. Так, из таблиц мы находим, что плотность воды при 5 , 10 и 20°C равна соответственно $0,999992 \text{ см}^{-3} \text{ г}$, $0,999727 \text{ см}^{-3} \text{ г}$ и $0,998230 \text{ см}^{-3} \text{ г}$, причём абсолютные ошибки этих величин не превышают одной единицы в последнем знаке, что отвечает указанному выше требованию. Таким образом, эти величины с учётом их абсолютных ошибок следовало бы написать так:

$$D_5 = (0,999992 \pm 0,000001) \text{ см}^{-3} \text{ г},$$

$$D_{10} = (0,999727 \pm 0,000001) \text{ см}^{-3} \text{ г},$$

$$D_{20} = (0,998230 \pm 0,000001) \text{ см}^{-3} \text{ г}.$$

В таких случаях абсолютные ошибки для каждого числа в отдельности обычно не указывают.

Пример 2. Точные измерения абсолютной величины ускорения силы тяжести дают возможность определять эту величину с точностью до $\pm 0,003 \text{ см сек}^{-2}$. Вследствие этого результаты одного из подобных измерений были даны в таком виде: $g = (980,274 \pm 0,003) \text{ см сек}^{-2}$. В этом

случае, очевидно, было бы совершенно неправильным не указать величину абсолютной ошибки.

Пример 3. Очень точное вычисление числа Авогадро N привело к выводу, что его наиболее вероятное значение таково:

$$N = (6,0230 \pm 0,0005) \cdot 10^{23} \text{ моль}^{-1}.$$

В этом случае мы, очевидно, во-первых, должны указать величину абсолютной ошибки и, во-вторых, не можем опустить нуль в конце первого числа, так как этот нуль является значащей цифрой, хотя и сомнительной. Если же значение числа Авогадро мы возьмём с точностью до второго десятичного знака, т. е. положим:

$$N = 6,02 \cdot 10^{23} \text{ моль}^{-1},$$

то все три знака в этом числе будут верными и его абсолютная ошибка не будет превышать 0,5 от одной единицы его последнего знака, как это можно непосредственно видеть.

Пример 4. Скорость света в пустоте c по последним данным, как уже было сказано выше, определяется числом

$$c = (2,99796 \pm 0,00004) \cdot 10^5 \text{ км сек}^{-1}.$$

Эту величину можно написать также в таком виде:

$$c = (299\,796 \pm 4) \text{ км сек}^{-1}.$$

Величину абсолютной ошибки необходимо, очевидно, указать, так как она достигает четырёх единиц в последнем знаке.

Здесь необходимо также обратить внимание на следующее:

Давая округлённое значение скорости света в пустоте, очень часто пишут его в таком виде:

$$c = 300\,000 \text{ км сек}^{-1}.$$

Это является совершенно недопустимым: последнее число показывает, во-первых, что все его знаки, за исключением последнего, являются верными и, во-вторых, что значение скорости света в пустоте нам известно с точностью до $\pm 1 \text{ км сек}^{-1}$, тогда как в действительности и то и другое является неверным. В подобных случаях необходимо сле-

довать общему правилу, т. е. указывать в приближённых числах только верные цифры, а также те, которые получаются в результате применения правила дополнения; при этом, если оказывается необходимым, вводят множитель в виде десяти в соответствующей степени. Так, значение скорости света в пустоте, если мы берем его с округлением, следует написать в таком виде:

$$c = 3 \cdot 10^5 \text{ км сек}^{-1}.$$

Такая форма записи приближённых значений является общепринятой; она применяется во всех случаях, где введение нулей могло бы повести к неправильному истолкованию ошибки приближённого числа.

Пример 5. Очень точное вычисление заряда e электрона привело к такой величине:

$$e = (4,8029 \pm 0,0005) \cdot 10^{-10} \text{ см}^{\frac{3}{2}} \text{ г}^{\frac{1}{2}} \text{ сек}^{-1}.$$

Для этой же величины иногда берут округлённое значение:

$$e = 4,803 \cdot 10^{-10} \text{ см}^{\frac{3}{2}} \text{ г}^{\frac{1}{2}} \text{ сек}^{-1}.$$

В первом случае мы, очевидно, должны указать величину абсолютной ошибки, во втором случае её можно не указывать.

В дальнейшем мы будем, как общее правило, приводя результаты измерений, указывать их ошибки, абсолютные или относительные. Если же в некоторых отдельных случаях величина ошибки не будет дана, то следует считать, что абсолютная ошибка данного числа равна ± 1 его последнего знака.

2) Определение относительной ошибки. Величина относительной ошибки приближённых чисел определяется формулой (4'):

$$\delta = \pm \frac{e}{A}. \quad (4')$$

Зная величину абсолютной ошибки приближённого значения числа, можно по этой формуле найти и его относительную ошибку. Но на практике для быстрого

определения относительных ошибок и упрощения арифметических вычислений иногда пользуются более простой, хотя и менее точной формулой. При этом мы несколько увеличиваем относительную ошибку, т. е. вычисляем относительную ошибку, которая заведомо больше ошибки, определяемой формулой (4'). Ошибки, которые удовлетворяют такому условию, принято называть особым термином — предельные ошибки. Таким образом, предельная относительная ошибка $\delta_{\text{пр}}$ должна удовлетворять неравенству

$$\delta_{\text{пр}} > \delta.$$

Для нахождения предельной относительной ошибки мы несколько уменьшаем число A и одновременно делаем его более простым. С этой целью все значащие цифры числа A , кроме первой, мы заменяем нулями. Так как число A может быть как целым числом, так и десятичной дробью, правильной или неправильной, то такую операцию в общем случае можно рассматривать как замену числа A некоторым числом N , определяемым выражением

$$N = k \cdot 10^{(n-1)-p},$$

где k обозначает первую значащую цифру числа A , а n и p обозначают соответственно число значащих цифр в нём и число десятичных знаков после запятой.

Таким образом, мы можем написать:

$$\delta_{\text{пр}} = \pm \frac{s}{k \cdot 10^{(n-1)-p}}. \quad (6)$$

На основании этой формулы мы можем сделать такой общий вывод:

Если первая значащая цифра некоторого приближённого числа A равна k и это число имеет n значащих цифр и p десятичных знаков, то его относительная ошибка меньше величины, определяемой формулой (6).

Рассмотрим несколько примеров.

Пример 1. Найдём предельную относительную ошибку одного из приведённых выше значений плотности воды при различных температурах. Например, для 10°C имеем:

$$D_{10} = (0,999727 \pm 0,000001) \text{ см}^{-3} \text{ г}.$$

В данном случае $n=6$, $p=6$. Применяя формулу (6), находим:

$$\delta_{\text{пр}} = \pm \frac{10^{-6}}{9 \cdot 10^{-1}} = \pm \frac{10^{-5}}{9} = \pm 0,0000011 = \pm 0,00011\%.$$

Более точное вычисление относительной ошибки по формуле (4') даёт:

$$\delta = \pm \frac{1}{999727} = \pm 0,000001 = \pm 0,0001\%,$$

т. е. величину несколько меньшую, хотя и очень близкую.

Пример 2. Найдём предельную относительную ошибку того значения абсолютной величины g , которое было приведено выше:

$$g = (980,274 \pm 0,003) \text{ см сек}^{-2}.$$

В этом случае $n=6$, $p=3$. По формуле (6) находим:

$$\delta_{\text{пр}} = \pm \frac{3 \cdot 10^{-3}}{9 \cdot 10^2} = \pm \frac{10^{-5}}{3} = \pm 0,0000033 = \pm 0,00033\%.$$

Более точное вычисление относительной ошибки по формуле (4') приводит к результату:

$$\delta = \pm \frac{3}{980274} = \pm 0,000003 = \pm 0,0003\%,$$

т. е. к результату, несколько меньшему, но близкому.

Пример 3. Найдём предельную относительную ошибку значения скорости света в пустоте c , которое даётся в таком виде:

$$c = (299\,796 \pm 4) \text{ км сек}^{-1}.$$

В данном случае $n=6$, $p=0$. Применяя формулу (6), находим:

$$\delta_{\text{пр}} = \pm \frac{4}{2 \cdot 10^5} = \pm 0,00002 = \pm 0,002\%.$$

Более точное значение относительной ошибки, вычисленное по формуле (4'), оказывается равным:

$$\delta = \pm \frac{4}{299\,796} = \pm 0,000013 = \pm 0,0013\%,$$

т. е. получается значительно меньшая величина. Причина большого расхождения между значениями $\delta_{\text{пр}}$ и δ в данном случае понятна: величина, которую мы поставили в знаменателе при вычислении $\delta_{\text{пр}}$, т. е. $2 \cdot 10^5 \text{ км сек}^{-1}$, очень сильно отличается от значения c , которое почти равно $3 \cdot 10^5 \text{ км сек}^{-1}$.

Пример 4. Найдём предельную относительную ошибку значения заряда электрона, данного числом

$$e := (4,8029 \pm 0,0005) \cdot 10^{-10} \text{ см}^{\frac{3}{2}} \text{ э}^{\frac{1}{2}} \text{ сек}^{-1}.$$

В данном случае $n=5$, $p=4$, и показатель степени в знаменателе формулы (6) оказывается равным нулю. По этой формуле находим:

$$\delta_{\text{пр}} = \pm \frac{5 \cdot 10^{-4}}{4} = \pm 0,000125 = \pm 0,0125\%.$$

Более точное вычисление по формуле (4') даёт:

$$\delta = \pm \frac{5}{48\,029} = \pm 0,000104 = \pm 0,01\%,$$

т. е. несколько меньшее значение, хотя и близкое.

Пример 5. Найдём предельную относительную ошибку приближённого значения π , взятого с пятью десятичными знаками:

$$\pi = 3,14159.$$

В данном случае $n=6$, $p=5$. Что же касается абсолютной ошибки, то для чисел этого типа её величина не может превышать 0,5 единицы последнего знака, т. е. в данном случае не может превышать $5 \cdot 10^{-6}$. Таким образом, по формуле (6) находим:

$$\delta_{\text{пр}} = \frac{5 \cdot 10^{-6}}{3} = 0,0000017 = 0,00017\%.$$

Более точное вычисление относительной ошибки этого значения даёт величину (стр. 25):

$$\delta = 0,00008\%,$$

т. е. значительно меньшую ошибку. Причина большого расхождения между $\delta_{\text{пр}}$ и δ в данном случае понятна: при вычислении $\delta_{\text{пр}}$ мы допустили, что абсолютная ошибка равна $0,5 \cdot 10^{-6}$, тогда как абсолютная ошибка этого приближения π почти в два раза меньше, как это можно видеть из таблицы приближённых значений π . Если из этой таблицы взять более точное значение абсолютной ошибки для данного значения π , то $\delta_{\text{пр}}$ оказывается равной

$$\delta_{\text{пр}} = \frac{2,65 \cdot 10^{-6}}{3} = 0,000088\%.$$

Это значение предельной относительной ошибки значительно ближе к значению δ , но вместе с тем $\delta_{\text{пр}}$ попрежнему больше δ .

Эти примеры показывают, что предельные относительные ошибки мы можем находить при помощи очень простых вычислений, и что значения предельных относительных ошибок в большинстве случаев оказываются достаточно близкими к более точным значениям относительных ошибок, которые вычисляются по формуле (4'). Условия, при которых $\delta_{\text{пр}}$ и δ могут значительно расходиться, становятся ясными из разобранных примеров.

3) Обратная задача, т. е. определение числа верных знаков приближённых чисел в зависимости от величины их относительных ошибок, вообще говоря, является несколько более сложной. Для элементарного решения этой задачи в простейших случаях достаточно определить абсолютную ошибку приближённого числа, величина которой даёт возможность непосредственно установить его верные знаки.

Из формулы (5'), выведенной раньше:

$$\epsilon = \pm A \delta, \quad (5')$$

мы знаем, что для определения абсолютной ошибки приближённого числа следует это число умножить на его относительную ошибку. Это вычисление абсолютных ошибок очень часто производят упрощённым приёмом, причём вычисляют ошибку, заведомо несколько большую ϵ , т. е. вычисляют предельную абсолютную

ошибку $\epsilon_{\text{пр}}$, которая удовлетворяет неравенству:

$$\epsilon_{\text{пр}} > \epsilon.$$

Для нахождения предельной абсолютной ошибки мы несколько увеличиваем число A и одновременно делаем его более простым. Для этого первую значащую цифру числа A мы увеличиваем на единицу, а остальные цифры заменяем нулями. Рассуждениями, вполне аналогичными тем, которые имели место при определении $\delta_{\text{пр}}$, мы приходим к выводу, что для этого число A следует заменить некоторым числом N , которое определяется выражением

$$N = (k + 1) \cdot 10^{(n-1)-p},$$

где попрежнему k обозначает первую значащую цифру числа A , а n и p обозначают число значащих цифр в нём и число десятичных знаков после запятой.

Таким образом, мы можем написать:

$$\epsilon_{\text{пр}} = (k + 1) 10^{(n-1)-p} \cdot \delta. \quad (7)$$

На основании этой формулы мы можем сделать такой общий вывод:

Если первая значащая цифра приближённого числа A равна k и число имеет n значащих цифр и p десятичных знаков, то его абсолютная ошибка меньше величины, определяемой формулой (7).

Рассмотрим несколько примеров.

Пример 1. Определим число верных знаков в табличном значении плотности ртути при 0°C :

$$d_0 = 13,5955 \text{ см}^{-3} \text{ г},$$

если известно, что относительная ошибка этой величины равна $0,0007\%$, или $7 \cdot 10^{-6}$.

В данном случае $n=6$, $p=4$. По формуле (7) находим:

$$\epsilon_{\text{пр}} = \pm 2 \cdot 10 \cdot 7 \cdot 10^{-6} = \pm 14 \cdot 10^{-5} = \pm 0,00014.$$

Отсюда мы видим, что в табличном значении плотности ртути при 0°C все цифры верны, за исключением последней, которая может отличаться приблизительно на величину ± 1 .

Более точное вычисление абсолютной ошибки по формуле (5') даёт такое значение ϵ :

$$\epsilon = \pm 13,5955 \cdot 7 \cdot 10^{-6} = \pm 0,000095 = \pm 0,0001.$$

Мы видим, что ϵ несколько меньше $\epsilon_{\text{пр}}$, но полученный нами вывод о числе верных знаков оказывается правильным.

Пример 2. Значение гравитационной постоянной κ принимается равным

$$\kappa = 6,67 \cdot 10^{-8} \text{ см}^3 \text{ г}^{-1} \text{ сек}^{-2}.$$

Многочисленные измерения этой величины выполнялись с различной точностью, среднюю ошибку измерений можно считать приближённо равной $0,8\%$, или $8 \cdot 10^{-3}$. Окончательный результат был вычислен с точностью до второго десятичного знака. Попробуем определить, является ли такая точность вычислений достаточной и не следует ли вычислить ещё третий десятичный знак.

Находим по формуле (7) предельную абсолютную ошибку этого числа, принимая во внимание, что в данном случае $n=3$, $p=2$:

$$\epsilon_{\text{пр}} = \pm 7 \cdot 0,008 = \pm 0,056.$$

Более точное вычисление абсолютной ошибки по формуле (5') приводит к такой величине:

$$\epsilon = \pm 6,67 \cdot 0,008 = \pm 0,05.$$

Мы видим, что ϵ оказывается несколько меньше $\epsilon_{\text{пр}}$. Одновременно величина абсолютной ошибки показывает, что в значении гравитационной постоянной второй десятичный знак является сомнительным, вследствие этого вычислять третий знак нет никаких оснований.

Пример 3. Скорость света в пустоте, измеренная с относительной ошибкой, равной $0,0014\%$, или $14 \cdot 10^{-6}$, оказалась равной $299\,796 \text{ км сек}^{-1}$. Определим верные знаки в этом числе.

В данном случае $n=6$, $p=0$. По формуле (7) находим:

$$\epsilon_{\text{пр}} = \pm 3 \cdot 10^5 \cdot 14 \cdot 10^{-6} = \pm 42 \cdot 10^{-1} = \pm 4,2.$$

Более точное вычисление по формуле (5') даёт:

$$\epsilon = \pm 299\,796 \cdot 14 \cdot 10^{-6} = \pm 4,19966.$$

Таким образом, хотя ϵ оказывается меньше $\epsilon_{\text{пр}}$, но их значения почти совпадают. Это объясняется тем, что величина $3 \cdot 10^5 \text{ км сек}^{-1}$, которой мы пользовались при определении $\epsilon_{\text{пр}}$, почти равна значению скорости света в пустоте.

Одновременно величина абсолютной ошибки показывает, что в числе $299\,796 \text{ км сек}^{-1}$ верными являются пять первых цифр, последняя цифра может отличаться на ± 4 единицы.

Пример 4. Определим верные знаки в значении постоянной Планка h :

$$h = 6,62 \cdot 10^{27} \text{ эрг сек},$$

которая определена с относительной ошибкой, равной 0,15%, или 0,0015.

В данном случае $n=3$, $p=2$; по формуле (7) находим:

$$\epsilon_{\text{пр}} = \pm 7 \cdot 0,0015 = \pm 0,0105.$$

Тот же результат получаем при более точном определении абсолютной ошибки по формуле (5'):

$$\epsilon = \pm 6,62 \cdot 0,0015 = \pm 0,00993 = \pm 0,01.$$

Попрежнему ϵ остаётся меньше $\epsilon_{\text{пр}}$, хотя оба эти значения очень близки. Величина абсолютной ошибки показывает, что в значении постоянной Планка второй десятичный знак является сомнительным.

Пример 5. Заряд e электрона, измеренный с относительной ошибкой, равной 0,01%, или 10^{-4} , оказывается равным:

$$e = 4,8029 \text{ см}^{\frac{3}{2}} \text{ г}^{\frac{1}{2}} \text{ сек}^{-1}.$$

Определим верные знаки в этом числе.

В данном случае $n=5$, $p=4$; по формуле (7) находим:

$$\epsilon_{\text{пр}} = \pm 5 \cdot 10^{-4} = \pm 0,0005.$$

Более точное вычисление абсолютной ошибки по формуле (5') даёт:

$$\epsilon = \pm 4,8029 \cdot 10^{-4} = \pm 0,00048 = \pm 0,0005.$$

Предельная абсолютная ошибка больше ϵ , но их значения очень близки. Величина абсолютной ошибки показывает, что в числе $4,8029 \text{ см}^{\frac{3}{2}} \text{ г}^{\frac{1}{2}} \text{ сек}^{-1}$ последняя цифра может отличаться на ± 5 единиц.

Пример 6. Одной из наиболее точно определённых физических констант является так называемая постоянная Ридберга. Её значение для водорода, определённое с относительной ошибкой, равной 0,00005%, или $5 \cdot 10^{-7}$, иногда пишут в таком виде:

$$R_{\text{H}} = 109\,677,691.$$

Определим верные знаки в этом числе.

В данном случае $n=9$, $p=3$. По формуле (7) находим:

$$\epsilon_{\text{пр}} = \pm 2 \cdot 10^5 \cdot 5 \cdot 10^{-7} = \pm 0,1.$$

Более точное определение абсолютной ошибки по формуле (5') даёт:

$$\epsilon = \pm 109\,677,691 \cdot 0,5 \cdot 10^{-7} = \pm 0,0548.$$

Большая разница в значениях ϵ и $\epsilon_{\text{пр}}$ в данном случае объясняется тем, что, увеличив на единицу первую цифру постоянной R_{H} , мы её значение почти удвоили. Что же касается верных цифр в числе R_{H} , то величина абсолютных ошибок показывает, что в этом числе первую цифру после запятой можно считать верной, но вторая цифра после запятой является безусловно сомнительной и может отличаться на ± 5 единиц. Таким образом, нет оснований писать это число с тремя десятичными знаками, как оно было написано выше, и при всех вычислениях, несколько не понижая их точности, можно R_{H} принимать равным 109 677,69, как это обычно и делают.

Вычисление предельных абсолютных ошибок, как показывают эти примеры, является очень простой операцией. Значения предельных абсолютных ошибок в большинстве случаев оказываются близкими к более точным значениям абсолютных ошибок, которые вычисляются по формуле (5'); это даёт возможность определять верные знаки в приближённых числах непосредственно по величине их предельных ошибок. Условия, при которых $\epsilon_{\text{пр}}$

и ϵ могут значительно расходиться, становятся ясными из разобранных примеров.

С вычислениями предельных ошибок, относительных и абсолютных, мы достаточно часто встречаемся на практике, в особенности в тех случаях, когда производим какие-либо сложные вычисления. В этих случаях на промежуточных этапах вычислений часто получаются результаты, которые имеют излишне большое число знаков, и для того, чтобы определить, сколько знаков можно отбросить, не нарушая общей точности вычислений, обычно приходится прибегать к вычислению предельных ошибок. С такими случаями нам придётся встретиться в дальнейшем.

ГЛАВА II

ОСНОВНЫЕ АРИФМЕТИЧЕСКИЕ ДЕЙСТВИЯ С ПРИБЛИЖЁННЫМИ ЗНАЧЕНИЯМИ ЧИСЕЛ

1. Малые величины различных порядков. Известно, что понятия большой и малой величин, больших и малых чисел являются понятиями вполне относительными. Например, дробь 10^{-3} велика по сравнению с дробью 10^{-6} , но мала по сравнению с единицей; единица велика по сравнению с дробью 10^{-3} , но мала по сравнению с числом 10^3 и т. д. Такой ряд последовательных сопоставлений различных чисел или величин можно, очевидно, продолжать сколь угодно далеко.

Несмотря на эту явную относительность понятий большого и малого, всё же оказывается возможным при оценке относительной величины различных объектов вносить некоторую определённую, хотя тоже только условную. Это достигается тем, что вводят понятие о степени или о порядке малых величин; иначе говоря, условно принимают, что малые величины можно делить на группы в зависимости от их относительных размеров. Так, дробь 10^{-6} , малая по сравнению с дробью 10^{-3} , мала также по сравнению с единицей и с числом 10^3 . Но единица велика по сравнению с дробью 10^{-3} , а число 10^3 велико по сравнению с единицей. Вследствие этого, если мы условно примем, что дробь 10^{-6} по сравнению с дробью 10^{-3} является малой величиной первого порядка, то можно считать, что та же дробь 10^{-6} является малой величиной второго порядка по сравнению с единицей и третьего порядка по сравнению с числом 10^3 .

Нетрудно видеть, что такое деление малых величин на различные группы в зависимости от их порядка является вполне условным, так как всегда можно ввести некоторые промежуточные величины, и это неизбежно должно вызывать изменение порядка остальных величин. Однако, несмотря на это, всё же оказалось, что деление малых величин на различные порядки дало возможность установить определённые правила арифметических действий с приближёнными числами и вывести ряд соответствующих формул, что в конечном результате значительно упростило технику вычислений.

Так как границы, отделяющие различные порядки малых величин, являются вполне условными, то в общем случае мы можем допустить, что если некоторая величина уменьшается в n раз, то она переходит в следующий порядок малых величин. Таким образом, если мы выберем некоторую величину N и по отношению к ней будем устанавливать порядок малых величин, то очевидно, что величины

$$N : n = \frac{N}{n}, \quad \frac{N}{n} : n = \frac{N}{n^2}, \quad \frac{N}{n^2} : n = \frac{N}{n^3}, \dots \quad (1)$$

следует считать малыми величинами соответственно первого, второго, третьего и т. д. порядков по отношению к величине N . Числовое значение n в этом ряду равенств для различных случаев мы можем брать совершенно различным; кроме того, n может изменяться в достаточно широких пределах даже для одного и того же случая. Так, при грубых предварительных подсчётах иногда можно принимать n приближённо равным от 100 до 200, т. е. если две величины отличаются одна от другой в 100—200 раз, то мы относим их в этих случаях к двум различным порядкам. Но при более точных вычислениях такое предположение становится непригодным, и обычно приходится предполагать, что n имеет значения в пределах приближённо от 1000 до 10 000, т. е. что только величины, меньшие данной в 1000—10 000 раз, можно относить к следующему порядку.

Пусть α является малой величиной первого порядка по отношению к единице, т. е. очень малой десятичной

дробью порядка 0,001—0,0001. На основании первого из равенств (1) имеем:

$$\alpha = \frac{1}{n}.$$

Возводя обе части этого равенства последовательно во вторую, третью и т. д. степени, находим:

$$\alpha^2 = \frac{1}{n^2}; \quad \alpha^3 = \frac{1}{n^3}; \dots$$

Из сравнения этих равенств с равенствами (1) следует, что:

Если малые величины первого порядка по отношению к единице возводить во вторую, третью и т. д. степени, то получаются малые величины соответственно второго, третьего и т. д. порядков по отношению к единице.

Если, далее, мы обозначим через $\alpha_1, \alpha_2, \alpha_3, \dots$ ряд малых величин первого порядка по отношению к единице, то из столь же элементарных соображений мы приходим к следующим выводам:

1) Произведения двух малых величин первого порядка, т. е. выражения вида

$$\alpha_1 \cdot \alpha_2, \alpha_2 \cdot \alpha_3, \alpha_1 \cdot \alpha_3, \dots,$$

следует считать малыми величинами второго порядка.

2) Произведения трёх малых величин первого порядка или произведения малой величины первого порядка на малую величину второго порядка, т. е. выражения вида

$$\alpha_1 \cdot \alpha_2 \cdot \alpha_3, \alpha_1 \cdot \alpha_2^2, \alpha_1 \cdot \alpha_3^2, \alpha_1^2 \cdot \alpha_2, \dots,$$

следует считать малыми величинами третьего порядка и т. д.

3) В соответствии с этим такие выражения, как например,

$$\frac{\alpha_1 \alpha_2}{\alpha_3}, \frac{\alpha_1 \alpha_3}{\alpha_2}, \frac{\alpha_2 \alpha_3}{\alpha_1}, \dots$$

или

$$\frac{\alpha_1 \cdot \alpha_2^2}{\alpha_3^2}, \frac{\alpha_1^2 \cdot \alpha_2}{\alpha_3^2}, \frac{\alpha_1 \cdot \alpha_3^2}{\alpha_2^2}, \dots,$$

остаются, очевидно, малыми величинами первого порядка.

Точно так же выражения

$$\frac{a_1 \cdot a_2^2}{a_3}, \frac{a_1 \cdot a_3^2}{a_2}, \frac{a_1 \cdot a_2^3}{a_3^2}, \frac{a_1 \cdot a_3^3}{a_2^2}, \dots$$

следует считать малыми величинами второго порядка, а такие выражения, как

$$\frac{a_1^2 \cdot a_2^2}{a_3}, \frac{a_1^3 \cdot a_3}{a_2}, \frac{a_1 \cdot a_2^3}{a_3}, \dots$$

являются, очевидно, малыми величинами третьего порядка.

Особое положение занимает частное от деления двух малых величин одного порядка, т. е. выражения вида

$$\frac{a_1}{a_2}, \frac{a_2}{a_3}, \frac{a_1^2}{a_2^2}, \frac{a_1 \cdot a_2}{a_3^2}, \frac{a_1^2 \cdot a_3}{a_2^3}.$$

Относительно величины этих дробей трудно сказать заранее что-либо определённое, так как малые величины, даже в пределах одного порядка, могут значительно различаться между собой; вследствие этого их отношение может быть и большой, и малой величиной, и величиной, близкой к единице.

Таким образом, можно сказать, что

Если имеются выражения, которые представляют собой произведения или степени малых величин, то установить их порядок не вызывает затруднений, но при определении величины отношения двух малых величин одинакового порядка следует соблюдать большую осторожность.

Далее нетрудно видеть, что сумму малых величин какого-либо одного порядка следует считать малой величиной того же порядка; исключения здесь мы могли бы встретить только в тех случаях, когда число слагаемых в этой сумме оказывается настолько большим, что начинает приближаться к числовому значению n , которое мы принимаем в данном случае. В этих случаях, на практике крайне редких, порядок суммы должен был бы, очевидно, понизиться на единицу. Однако для окончательного реше-

ния этого вопроса необходимо, чтобы были известны знаки слагаемых этой суммы.

2. Формулы для приближённых вычислений. При всех вычислениях с приближёнными числами принято применять следующее общее правило.

Все арифметические действия с приближёнными числами производят с точностью до малых величин только определённого порядка, в большинстве случаев с точностью до малых величин того же порядка, который мы имеем в начальных данных; что же касается малых величин более высоких порядков, то их при вычислениях во внимание не принимают. Так, если $a_1, a_2, a_3, \dots$ являются малыми величинами первого порядка по отношению к данным нам величинам, то все вычисления с последними следует производить с точностью до малых величин тоже только первого порядка, а все малые величины второго и более высоких порядков, если они получаются при вычислениях, не следует принимать во внимание, отбрасывая их в окончательном результате.

Это правило значительно упрощает всю технику вычислений и одновременно позволяет вывести ряд очень простых приближённых формул, которые дают достаточно точные результаты и вследствие этого находят весьма широкие применения. Так, например, если a является малой величиной первого порядка по отношению к единице, то согласно приближённым формулам можно принимать:

$$\begin{aligned} (1 \pm a)^2 &= 1 \pm 2a, \\ (1 \pm a)^3 &= 1 \pm 3a. \end{aligned} \quad (2)$$

Точно так же, если $a_1, a_2, a_3, \dots$ обозначают попрежнему малые величины первого порядка по отношению к единице, то согласно приближённой формуле можно принимать:

$$(1 \pm a_1)(1 \pm a_2)(1 \pm a_3) \dots = 1 \pm a_1 \pm a_2 \pm a_3 \pm \dots \quad (3)$$

Все подобные формулы, весьма разнообразные (стр. 373), очень просто выводятся математически. Так, для первой

из формул (2) имеем:

$$(1 \pm \alpha)^2 = 1 \pm 2\alpha + \alpha^2.$$

Отбрасывая в этой формуле α^2 как малую величину второго порядка по отношению к единице, получаем приближённую формулу (2).

В случае формулы (3), ограничиваясь тремя множителями, имеем:

$$(1 \pm \alpha_1)(1 \pm \alpha_2)(1 \pm \alpha_3) = \\ = 1 \pm \alpha_1 \pm \alpha_2 \pm \alpha_3 \pm \alpha_1\alpha_2 \pm \alpha_1\alpha_3 \pm \alpha_2\alpha_3 \pm \alpha_1\alpha_2\alpha_3.$$

В этой формуле четыре последних члена мы отбрасываем как малые величины второго и третьего порядка по отношению к единице и получаем в результате приближённую формулу (3).

Во всех разделах физики очень широко применяются приближённые формулы чрезвычайно разнообразных видов. Подобными формулами можно всегда пользоваться при вычислениях, но только при одном неперемennom условии: применяя ту или иную приближённую формулу, мы не должны выходить за пределы тех ограничений, которые были приняты при выводе этой формулы. Для того чтобы лучше разобраться в этих вопросах, рассмотрим несколько примеров.

Пример 1. Точная формула, которой в общем случае определяется период колебания T математического маятника, имеет такой вид:

$$T = 2\pi \sqrt{\frac{l}{g}} \left[1 + \left(\frac{1}{2}\right)^2 \sin^2 \frac{\alpha}{2} + \left(\frac{1 \cdot 3}{2 \cdot 4}\right)^2 \sin^4 \frac{\alpha}{2} + \right. \\ \left. + \left(\frac{1 \cdot 3 \cdot 5}{2 \cdot 4 \cdot 6}\right)^2 \sin^6 \frac{\alpha}{2} + \dots \right], \quad (4)$$

где l и g обозначают длину маятника и ускорение силы тяжести, а α обозначает амплитуду колебаний, т. е. угол отклонения маятника от вертикали.

С другой стороны, очень часто период колебания математического маятника вычисляют, пользуясь приближённой формулой

$$T = 2\pi \sqrt{\frac{l}{g}}, \quad (4')$$

получаемой из формулы (4), если мы в её правой части отбросим в скобках все члены, содержащие знак синуса.

Принято считать, что формулой (4') можно пользоваться при малых амплитудах колебаний, однако при этом величина амплитуды более точно обыкновенно не указывается. Вследствие этого пределы возможного применения формулы (4') остаются неясными; точно так же остаётся неясным, при каких условиях необходимо переходить к формуле (4), и сколько членов бесконечного ряда в этой формуле следует принимать во внимание. В эти вопросы можно внести некоторую ясность, пользуясь соображениями, изложенными выше, хотя при этом, конечно, необходимо иметь в виду, что здесь многое попрежнему определяется той точностью, с которой мы производим вычисления.

Рассматривая бесконечный ряд в правой части формулы (4), мы видим, что при малых амплитудах, когда $\sin \frac{\alpha}{2}$ очень малая дробь, близкая к нулю, все члены этого ряда, начиная со второго, являются малыми величинами по сравнению с единицей, причём мы можем считать, что их порядок с каждым членом возрастает на две единицы. Таким образом, мы имеем ряд с очень быстро убывающими по величине членами. Вследствие этого даже при очень точных вычислениях мы можем ограничиваться только вторым членом ряда, т. е. пользоваться формулой вида:

$$T = 2\pi \sqrt{\frac{l}{g}} \left[1 + \left(\frac{1}{2}\right)^2 \sin^2 \frac{\alpha}{2} \right]. \quad (4'')$$

Если от этой формулы мы переходим к формуле (4'), то, очевидно, мы полагаем, что второй член в скобках является очень малой величиной по сравнению с единицей. Вследствие этого при вычислениях обычной точности мы можем положить, например, что

$$\frac{1}{4} \sin^2 \frac{\alpha}{2} < 0,0001.$$

Из последнего выражения имеем:

$$\sin \frac{\alpha}{2} < 0,02.$$

Отсюда по таблицам находим, что α можно считать приближённо равным 2° .

Таким образом, мы видим, что при вычислениях обычной точности, когда относительная ошибка имеет величину около 0,01%, можно пользоваться формулой (4'), если амплитуда колебаний не выходит за пределы 2° . При больших амплитудах или при более точных вычислениях следует переходить к формуле (4''). Эту формулу можно считать очень точной, так что только в отдельных чрезвычайно редких случаях, например, при очень больших амплитудах, могла бы возникнуть необходимость ввести при вычислениях ещё следующий член ряда, содержащий $\sin \frac{\alpha}{2}$ в четвёртой степени.

Пример 2. Как известно, между термическим коэффициентом линейного расширения α твёрдых тел и их коэффициентом объёмного расширения β установлено простое соотношение

$$\beta = 3\alpha. \quad (5)$$

В таблицах обыкновенно даются значения α , и если необходимо найти значение коэффициента объёмного расширения какого-либо тела, то соответствующее табличное значение α утраивают.

Формула (5) является приближённой, и более точное соотношение между β и α , которое очень просто выводится алгебраически, имеет такой вид:

$$\beta = 3\alpha + 3\alpha^2 + \alpha^3. \quad (5')$$

Из этой формулы получается приближённая формула (5), если мы допускаем, что α является очень малой дробью, т. е. что

$$\alpha \ll 1.$$

Это действительно имеет место, например, по отношению к металлам, их сплавам, к твёрдым телам вообще, но вместе с тем из таблиц мы видим, что значения коэффициентов линейного расширения различных твёрдых тел очень сильно отличаются между собой, в сотни и тысячи раз. Вследствие этого может возникнуть вопрос, всегда ли

мы вправе пользоваться приближённой формулой (5) и не следует ли иногда переходить к более точной формуле (5'), например, в тех случаях, когда α имеет относительно большую величину.

Для того чтобы выяснить этот вопрос, попробуем найти значения коэффициента объёмного расширения, применяя последовательно обе формулы, приближённую (5) и более точную (5'). Возьмём какой-либо металл с относительно большой величиной коэффициента линейного расширения, например, алюминий. Из таблиц находим, что его линейный коэффициент расширения α_{Al} , средний для температурного интервала 0°C и 100°C , определяется такой величиной:

$$\alpha_{Al} = (0,238 \pm 0,001) \cdot 10^{-4} \text{ град}^{-1}.$$

Таким образом, в этом числе, как показывает его абсолютная ошибка, первые шесть десятичных знаков являются верными, а последний, седьмой, знак сомнительным.

Применяя формулу (5), находим непосредственно:

$$\beta_{Al} = (0,714 \pm 0,003) \cdot 10^{-4} \text{ град}^{-1},$$

причём в этом числе попрежнему первые шесть знаков мы должны считать верными, а последний, седьмой, сомнительным.

Если же значение β_{Al} мы будем вычислять, пользуясь более точной формулой (5'), то нам придётся проделать ряд сложных арифметических действий, так как значение α_{Al} придётся возводить во вторую и третью степень. В результате мы получим число, которое должно иметь двадцать один десятичный знак, но из этого числа десятичных знаков мы можем оставить только первые семь, а именно: шесть первых знаков верных и седьмой сомнительный; все остальные знаки в этом числе, начиная с восьмого, мы обязаны отбросить как неверные. Первые двенадцать десятичных знаков этого числа таковы:

$$0,71401699 \cdot 10^{-4}.$$

Ограничиваясь первыми семью десятичными знаками и вводя значение абсолютной ошибки, получаем:

$$\beta_{Al} = (0,714 \pm 0,003) \cdot 10^{-4} \text{ град}^{-1},$$

Таким образом, формула (5') приводит в точности к тому же значению коэффициента объёмного расширения алюминия, которое мы получили по приближённой формуле (5) простым умножением α_{A1} на три.

Пример 3. Ускорение силы тяжести уменьшается при вертикальном подъёме над поверхностью Земли по определённому закону; его иногда выражают приближённой формулой такого вида:

$$g_h = g_0 \left(1 - 2 \frac{h}{R} \right). \quad (6)$$

В этой формуле h и R обозначают высоту поднятия, считая её от уровня моря, и средний радиус земного шара, а g_0 и g_h обозначают ускорение силы тяжести соответственно на уровне моря и на высоте h .

Формула (6) выведена на основании закона всемирного тяготения, но при одном вполне определённом ограничении: высота поднятия h должна быть малой по сравнению с радиусом земли R :

$$h \ll R. \quad (7)$$

Более точное соотношение между g_h и g_0 , свободное от этого ограничения, имеет такой вид:

$$g_h = g_0 \cdot \frac{1}{1 + 2 \frac{h}{R} + \frac{h^2}{R^2}}. \quad (6')$$

Найдём значения g_h , вычисляя их по формулам (6) и (6') для одних и тех же высот, например для высот 1 км, 10 км и 20 км. Результаты этих вычислений с точностью до второго десятичного знака приведены в таблице I, где указана также относительная ошибка значений g_h , вычисленных по приближённой формуле (6).

Таблица I

h	1 км	10 км	20 км
Формула (6)	980,69	977,92	974,84
Формула (6')	980,69	977,93	974,87
δ	0%	0,001%	0,003%

Мы видим, что результаты, которые даёт приближённая формула (6), очень близки к более точным значениям g_h , которые мы находим по формуле (6'). Вследствие этого, если мы производим вычисления не особенно большой точности, например, с относительной ошибкой около 0,01%, то, очевидно, мы можем применять приближённую формулу (6) в широком интервале высот, до 20 км и даже несколько дальше. Вместе с тем, однако, следует заметить, что при дальнейшем значительном увеличении h формула (6) начинает давать неверные результаты; так, при высоте 100 км относительная ошибка, даваемая этой формулой, достигает 0,07%, а при высоте 500 км ошибка превышает 2%, что объясняется тем, что при таких значениях h основное условие (7) оказывается нарушенным. А если бы мы, забыв об этом условии, применили формулу (6) к очень большим высотам, например, положили бы h равным половине земного радиуса, то получился бы совершенно неправдоподобный результат: g оказалось бы равным нулю, а затем, при дальнейшем увеличении h , g получило бы отрицательное значение.

Рассмотренные примеры показывают, что приближённые формулы оказываются чрезвычайно полезными; они дают надёжные результаты и значительно сокращают вычислительные операции. Но каждую приближённую формулу необходимо предварительно исследовать, установить пределы её применимости, если они определены недостаточно точно, и не выходить далеко за эти пределы, так как в противном случае, как показывает последний пример, мы можем встретиться с совершенно неверными результатами.

3. Сложение и вычитание приближённых чисел. Производя вычисления с приближёнными числами, мы выполняем над ними различные арифметические действия, каждое из которых оказывает своё влияние на величину ошибки окончательного результата вычислений. Вследствие этого необходимо, очевидно, исследовать вопрос о том, какими ошибками должно сопровождаться в отдельности каждое из основных арифметических действий с приближёнными числами, т. е. сложение и вычитание, умноже-

ние и деление, возведение в степень и извлечение корня. Иными словами, необходимо найти общие выражения, которые определяют величину ошибок, абсолютных и относительных: 1) суммы и разности приближённых чисел, 2) их произведения и частного и 3) их степени и корня из них. Все эти вопросы решаются при помощи общего метода; он состоит в том, что мы берём приближённые числа в общем виде, выполняем над ними определённое арифметическое действие и, анализируя полученный результат, определяем величину его ошибок, абсолютной и относительной. Рассмотрим эти вопросы в указанном выше порядке.

1. *Ошибки суммы.* Возьмём несколько приближённых величин, выражая их в общей форме согласно уравнению (1') (стр. 27):

$$\begin{aligned} X_1 &= A_1 \pm \epsilon_1, \\ X_2 &= A_2 \pm \epsilon_2, \\ X_3 &= A_3 \pm \epsilon_3, \\ &\dots \dots \dots \\ &\dots \dots \dots \end{aligned}$$

Складывая эти уравнения, находим:

$$\begin{aligned} X_1 + X_2 + X_3 + \dots &= \\ &= A_1 + A_2 + A_3 + \dots \pm \epsilon_1 \pm \epsilon_2 \pm \epsilon_3 \pm \dots \end{aligned} \quad (8)$$

Из этого выражения мы видим, что сумма величин A и сумма ϵ служат соответственно приближённым значением суммы чисел X и её абсолютной ошибкой.

1) *Абсолютная ошибка суммы.* Из выражения (8) мы видим, что абсолютная ошибка суммы приближённых значений величин равняется сумме абсолютных ошибок всех слагаемых. Однако этот вывод не даёт нам возможности в общем случае определить величину абсолютной ошибки суммы, так как знаки абсолютных ошибок слагаемых в общем случае остаются нам неизвестными. Вследствие этого при сложении приближённых значений величин в тех случаях, когда знаки их абсолютных ошибок остаются неизвестными, всегда предпола-

гается, что имеет место наименее благоприятный случай, т. е. что абсолютные ошибки всех слагаемых имеют один и тот же знак, положительный или отрицательный. При таком предположении мы, возможно, преувеличиваем значение абсолютной ошибки суммы, но вместе с тем мы можем утверждать, что даже в самом неблагоприятном случае действительная ошибка суммы приближённых значений величин не будет превышать той, которую мы принимаем. Такую абсолютную ошибку мы вновь можем назвать предельной абсолютной ошибкой (стр. 41—42). Таким образом, обозначая предельную абсолютную ошибку суммы нескольких приближённых величин через $\epsilon_{\text{пр}}$, мы можем написать:

$$\epsilon_{\text{пр}} = \pm (|\epsilon_1| + |\epsilon_2| + |\epsilon_3| + \dots). \quad (9)$$

Предельная абсолютная ошибка суммы нескольких приближённых значений величин равняется сумме абсолютных значений абсолютных ошибок всех слагаемых.

Отсюда следует, что если среди приближённых значений величин, которые мы складываем, хотя бы одна величина имеет абсолютную ошибку, значительно бóльшую, чем ошибки остальных слагаемых, то и сумма этих приближённых значений величин будет иметь бóльшую абсолютную ошибку. Это даёт возможность при сложении приближённых значений величин вычислять сумму с округлением, заранее отбрасывая у некоторых слагаемых лишние знаки.

2) *Относительная ошибка суммы.* На основании общего определения относительной ошибки (формула (4'), стр. 37) находим из формулы (8), что относительная ошибка суммы нескольких приближённых величин определяется выражением

$$\delta = \frac{\pm \epsilon_1 \pm \epsilon_2 \pm \epsilon_3 \pm \dots}{A_1 + A_2 + A_3 + \dots}.$$

Это выражение показывает, что в общем случае мы вновь встречаемся с двойным знаком абсолютных ошибок

и не можем вследствие этого более точно определить значение δ . Поэтому находят величину предельной относительной ошибки, которая, очевидно, определяется выражением

$$\delta_{\text{пр}} = \pm \frac{\varepsilon_{\text{пр}}}{A_1 + A_2 + A_3 + \dots} = \pm \frac{|\varepsilon_1| + |\varepsilon_2| + |\varepsilon_3| + \dots}{A_1 + A_2 + A_3 + \dots}. \quad (10)$$

Из этого выражения непосредственно не видно, как зависит относительная ошибка суммы от относительных ошибок слагаемых, но этот вопрос нетрудно выяснить на основании следующих соображений.

В правой части формулы (10) абсолютные ошибки ε заменим произведениями приближённых чисел A на их относительные ошибки. Находим

$$\delta_{\text{пр}} = \pm \frac{A_1 |\delta_1| + A_2 |\delta_2| + A_3 |\delta_3| + \dots}{A_1 + A_2 + A_3 + \dots}.$$

В этом выражении среди относительных ошибок δ выберем наибольшую и наименьшую и обозначим их соответственно $\delta_{\text{макс}}$ и $\delta_{\text{мин}}$. Если в последнем выражении заменить все δ сначала $\delta_{\text{макс}}$, а затем $\delta_{\text{мин}}$, то, очевидно, можно написать:

$$\frac{A_1 |\delta_{\text{макс}}| + A_2 |\delta_{\text{макс}}| + A_3 |\delta_{\text{макс}}| + \dots}{A_1 + A_2 + A_3 + \dots} > \delta_{\text{пр}},$$

$$\delta_{\text{пр}} > \frac{A_1 |\delta_{\text{мин}}| + A_2 |\delta_{\text{мин}}| + A_3 |\delta_{\text{мин}}| + \dots}{A_1 + A_2 + A_3 + \dots}.$$

Отсюда имеем:

$$\delta_{\text{макс}} > \delta_{\text{пр}} > \delta_{\text{мин}}.$$

Таким образом, мы видим, что:

Предельная относительная ошибка суммы лежит в границах между наибольшей и наименьшей относительными ошибками слагаемых.

Очевидно, мы могли бы считать, что $\delta_{\text{пр}}$ равна $\delta_{\text{макс}}$, но во многих случаях это было бы излишним преувеличением ошибки суммы. Вследствие этого последние неравенства и наш вывод следует оставить в силе. Необходимо также указать, что здесь возможен ещё один частный

случай, на который следует обратить внимание. Мы разберём этот случай, ограничиваясь двумя слагаемыми A_1 и A_2 . При этом условии формула (10) получает вид:

$$\delta_{\text{пр}} = \pm \frac{|\varepsilon_1| + |\varepsilon_2|}{A_1 + A_2}.$$

Если в частном случае одно из чисел A , например A_2 , является точным, то его абсолютная ошибка ε_2 равна нулю, и из последней формулы мы находим:

$$\delta_{\text{пр}} = \pm \frac{\varepsilon_1}{A_1 + A_2}.$$

С другой стороны, на основании общего определения относительной ошибки можно написать:

$$\delta_1 = \pm \frac{\varepsilon_1}{A_1}.$$

Из сопоставления двух последних формул мы видим, что

$$\delta_{\text{пр}} < \delta_1,$$

т. е. что

При сложении приближённого числа с точным предельная относительная ошибка суммы оказывается меньше относительной ошибки приближённого числа, причём очевидно, что $\delta_{\text{пр}}$ должно быть тем меньше по сравнению с δ_1 , чем больше точное число A_2 по сравнению с приближённым числом A_1 . На это обстоятельство следует обратить внимание, так как с ним часто приходится встречаться на практике.

На практике при определении ошибок суммы ряда приближённых значений величин сначала находят предельную абсолютную ошибку суммы, а затем её предельную относительную ошибку, пользуясь соответственно формулами (9) и (10). Рассмотрим какой-нибудь пример.

Пример. Точное измерение сопротивлений r_1 , r_2 , r_3 и r_4 четырёх технических эталонов сопротивления на 100 ом, 250 ом, 1000 ом и 10 000 ом дало такие

результаты:

$$\begin{aligned} r_1 &= 100,12 \pm 0,01 \text{ ом} & (0,01\%), \\ r_2 &= 249,61 \pm 0,02 \text{ ом} & (0,008\% \approx 0,01\%), \\ r_3 &= 1001,2 \pm 0,1 \text{ ом} & (0,01\%), \\ r_4 &= 10\,003 \pm 1 \text{ ом} & (0,01\%). \end{aligned}$$

Относительные ошибки этих измерений, вычисленные по формуле (4') главы I, указаны в скобках. Определим общее сопротивление R при последовательном соединении всех этих сопротивлений:

$$R = r_1 + r_2 + r_3 + r_4.$$

Складываем величины r и находим по формуле (9) предельную абсолютную ошибку $\epsilon_{\text{пр}}$:

100,12	0,01
249,61	0,02
1001,20	0,1
10 003,00	1
$R = 11\,353,93$	$\epsilon_{\text{пр}} = 1,13$

Величина предельной абсолютной ошибки $\epsilon_{\text{пр}}$, приближённо равной $\pm 1 \text{ ом}$, показывает, что в значении R цифры, стоящие после запятой, следует отбросить как неверные. Вследствие этого окончательный результат мы должны написать в таком виде:

$$R = (11\,354 \pm 1) \text{ ом}.$$

Мы видим, что абсолютная ошибка суммы $\epsilon_{\text{пр}}$ в данном случае определяется исключительно абсолютной ошибкой одного из слагаемых, именно r_4 : его абсолютная ошибка значительно превосходит абсолютные ошибки остальных слагаемых, которые вследствие этого не оказывают заметного влияния на абсолютную ошибку суммы. Отсюда следует также, что сложение величин r можно выполнить несколько проще, не принимая во внимание последний знак r_1 и r_2 , так как этот знак в окончательном результате мы, очевидно, должны отбросить.

Вычисляя предельную относительную ошибку суммы по формуле (10), находим:

$$\delta_{\text{пр}} = \pm \frac{1,13}{11\,353,93} = \pm 0,009\% = \pm 0,01\%.$$

В данном случае величина относительной ошибки всех слагаемых одинакова, и их сумма сохраняет то же значение относительной ошибки.

II. *Ошибки разности.* Два приближённые числа, данные в общем виде:

$$\begin{aligned} X_1 &= A_1 \pm \epsilon_1, \\ X_2 &= A_2 \pm \epsilon_2, \end{aligned}$$

вычитаем одно из другого.

Находим:

$$X_1 - X_2 = A_1 - A_2 + (\pm \epsilon_1 \mp \epsilon_2).$$

1) *Абсолютная ошибка разности.* Величина, стоящая в скобках в правой части этого выражения, определяет абсолютную ошибку разности двух приближённых величин. Мы вновь видим, что определить эту величину более точно в общем случае невозможно вследствие двойного знака величин ϵ_1 и ϵ_2 . Поэтому опять предполагают наименее благоприятный случай, когда знаки ϵ_1 и ϵ_2 оказываются одинаковыми и находят предельную абсолютную ошибку $\epsilon_{\text{пр}}$ разности, которая, очевидно, определяется выражением

$$\epsilon_{\text{пр}} = \pm (|\epsilon_1| + |\epsilon_2|). \quad (11)$$

Таким образом, мы можем сказать, что:

Предельная абсолютная ошибка разности двух приближённых значений величин равняется сумме абсолютных значений абсолютных ошибок уменьшаемого и вычитаемого.

Это положение и то, которое мы имели раньше относительно предельной абсолютной ошибки суммы, можно, очевидно, объединить одной общей теоремой:

Предельная абсолютная ошибка алгебраической суммы нескольких приближённых

ных величин равняется сумме абсолютных значений абсолютных ошибок слагаемых.

2) Относительная ошибка разности. Из предыдущего нетрудно видеть, что относительная ошибка разности двух приближённых величин определяется выражением:

$$\delta = \frac{\pm \varepsilon_1 \pm \varepsilon_2}{A_1 - A_2}.$$

Величину δ мы вновь заменяем предельной относительной ошибкой $\delta_{\text{пр}}$, равной

$$\delta_{\text{пр}} = \pm \frac{\varepsilon_{\text{пр}}}{A_1 - A_2} = \pm \frac{|\varepsilon_1| + |\varepsilon_2|}{A_1 - A_2}. \quad (12)$$

Ошибки разности приближённых значений величин находятся в той же последовательности, как и ошибки суммы; т. е. сначала определяют предельную абсолютную ошибку по формуле (11), затем предельную относительную ошибку по формуле (12).

Несмотря на большую простоту последней формулы, всё же по отношению к предельной относительной ошибке разности необходимо сделать следующее очень существенное замечание:

Числитель в правой части формулы (12) является малой величиной по сравнению с A_1 и A_2 , что же касается знаменателя, то

1) знаменатель $A_1 - A_2$ является большой величиной, если A_1 много больше A_2 — в этом случае предельная относительная ошибка остаётся малой,

2) но если A_1 и A_2 близки друг к другу, то знаменатель $A_1 - A_2$ может оказаться малой величиной, а в таком случае $\delta_{\text{пр}}$ может получить большое значение.

Таким образом, мы видим, что:

Предельная относительная ошибка разности двух приближённых чисел, близких по величине, может оказаться большой величиной, даже если относительные ошибки каждого из этих чисел в отдельности являются малыми.

С этим обстоятельством часто приходится встречаться в практике измерений, поэтому разберём его на примерах.

Пример 1. При определении веса таких тел, как, например, очень маленькие капельки ртути, часто применяют следующий приём: на чувствительных аналитических весах очень тщательно взвешивают капельку ртути, помещая её на часовое стёклышко, вес которого предварительно столь же тщательно определяется; из этих двух взвешиваний находят затем вес капельки ртути. При одной из таких операций были получены следующие результаты:

Вес часового стекла	$(6102,37 \pm 0,05) \text{ мГ } (\delta = \pm 0,001\%)$
То же с каплей ртути	$(6109,21 \pm 0,05) \text{ мГ } (\delta = \pm 0,001\%)$
Вес капли ртути	$(6,84 \pm 0,1) \text{ мГ}$

Относительные ошибки показывают, что взвешивание было выполнено с высокой степенью точности, но, несмотря на это, окончательный результат наших измерений оказывается очень неточным. Действительно,

$$\delta_{\text{пр}} = \pm \frac{0,05 + 0,05}{6,84} = \pm 0,0146 = \pm 1,5\%.$$

Такой результат показывает, что аналитические весы, даже весьма чувствительные, мало пригодны для точного взвешивания очень лёгких тел. Вследствие этого были выработаны особые типы весов специально для очень малых тел весом в несколько миллиграммов, так называемые микровесы различных видов.

Пример 2. Ёмкость шарового конденсатора (рис. 1) определяется формулой

$$C = \frac{\varepsilon \cdot Rr}{R - r},$$

где R и r обозначают радиусы соответственно внешней и внутренней шаровых поверхностей конденсатора, а ε обозначает диэлектрическую постоянную среды, заполняющей промежуток между ними.

При сборке одного из таких конденсаторов величины R и r были тщательно измерены. Результаты измерений оказались таковы:

$$R = (5,27 \pm 0,01) \text{ см} \quad (\delta = \pm 0,2\%),$$

$$r = (5,012 \pm 0,002) \text{ см} \quad (\delta = \pm 0,04\%).$$

Определяя отсюда величину $R - r$, находим по формуле (12) её предельную ошибку:

$$\delta_{\text{пр}} = \pm \frac{0,01 + 0,002}{0,258} = \pm 0,05 = \pm 5\%.$$

Таким образом, относительная ошибка величины $R - r$ может достигать 5% при сравнительно большой точности измерений величин R и r в отдельности. Отсюда следует, что определение ёмкости шарового конденсатора этим приёмом может дать грубо неточные результаты.

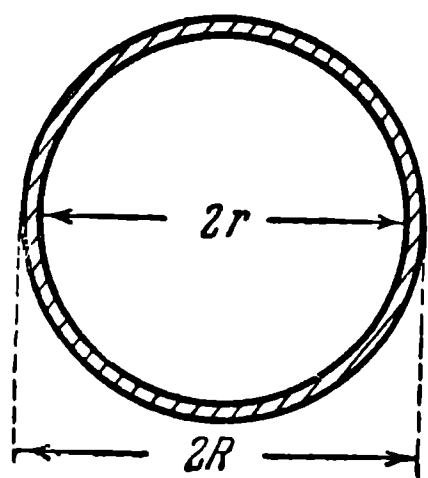


Рис. 1.

Такие результаты, на первый взгляд несколько неожиданные, мы можем встретить во всех случаях, когда в формулах оказывается в знаменателе разность двух близких между собой приближённых величин. Подобные формулы приходится тщательно исследовать. Иногда разность в знаменателе удаётся исключить при помощи математических преобразований, если же это оказывается невозможным, то часто приходится отказываться от такого метода и переходить к какому-либо другому методу измерений, свободному от таких осложнений.

4. Умножение и деление приближённых чисел.
I. *Ошибки произведения.* Несколько приближённых чисел, взятых в общем виде:

$$X_1 = A_1 \pm \epsilon_1, \quad X_2 = A_2 \pm \epsilon_2, \quad X_3 = A_3 \pm \epsilon_3, \dots$$

после умножения дают:

$$X_1 X_2 X_3 \dots = (A_1 \pm \epsilon_1)(A_2 \pm \epsilon_2)(A_3 \pm \epsilon_3) \dots \quad (13)$$

Из этой формулы можно определить величину ошибок произведения двух, трёх или в общем случае нескольких приближённых чисел, причём сначала удобнее найти их относительные ошибки.

1) Относительная ошибка произведения. Переходя в формуле (13) к относительным ошибкам и ограничиваясь двумя приближёнными числами, имеем:

$$X_1 X_2 = A_1(1 \pm \delta_1) \cdot A_2(1 \pm \delta_2) = A_1 A_2 (1 \pm \delta_1 \pm \delta_2 \pm \delta_1 \delta_2).$$

Произведение $\delta_1 \delta_2$ в правой части этой формулы мы можем отбросить как малую величину второго порядка по сравнению с единицей. Получаем:

$$X_1 X_2 = A_1 A_2 (1 \pm \delta_1 \pm \delta_2). \quad (14)$$

Совершенно так же для трёх приближённых множителей находим из формулы (13):

$$X_1 X_2 X_3 = A_1 A_2 A_3 (1 \pm \delta_1)(1 \pm \delta_2)(1 \pm \delta_3) =$$

$$= A_1 A_2 A_3 (1 \pm \delta_1 \pm \delta_2 \pm \delta_3 \pm \delta_1 \delta_2 \pm \delta_2 \delta_3 \pm \delta_1 \delta_3 \pm \delta_1 \delta_2 \delta_3).$$

Отбрасывая в скобках малые величины второго и третьего порядков по сравнению с единицей, находим:

$$X_1 X_2 X_3 = A_1 A_2 A_3 (1 \pm \delta_1 \pm \delta_2 \pm \delta_3). \quad (14')$$

Наконец, в общем случае нескольких приближённых множителей в результате таких же преобразований получаем выражение

$$X_1 X_2 X_3 \dots = A_1 A_2 A_3 \dots (1 \pm \delta_1 \pm \delta_2 \pm \delta_3 \pm \dots). \quad (14'')$$

Члены, которые стоят после единицы в скобках в правой части формул (14), (14') и (14''), представляют собой, очевидно, относительные ошибки произведения соответственно двух, трёх и нескольких множителей. Мы видим, что величина относительной ошибки произведения в общем виде определяется выражением

$$\delta = \pm \delta_1 \pm \delta_2 \pm \delta_3 \pm \dots$$

Наличие двойного знака при относительных ошибках отдельных множителей в правой части этой формулы заставляет нас переходить попрежнему к вычислению

предельной относительной ошибки $\delta_{\text{пр}}$, величина которой в общем случае нескольких множителей определяется, очевидно, выражением

$$\delta_{\text{пр}} = \pm (|\delta_1| + |\delta_2| + |\delta_3| + \dots). \quad (15)$$

Таким образом, мы приходим к выводу, что:

Предельная относительная ошибка произведения нескольких приближённых чисел равняется сумме абсолютных значений относительных ошибок всех множителей.

2) Абсолютная ошибка произведения. Из формулы (13), ограничиваясь двумя множителями, находим:

$$X_1 X_2 = A_1 A_2 \pm \epsilon_1 A_2 \pm \epsilon_2 A_1 \pm \epsilon_1 \epsilon_2.$$

По отношению к величине $A_1 A_2$ произведение $\epsilon_1 \epsilon_2$ является малой величиной второго порядка, которую можно отбросить, поэтому имеем:

$$X_1 X_2 = A_1 A_2 \pm \epsilon_1 A_2 \pm \epsilon_2 A_1.$$

Абсолютной ошибкой ϵ этого произведения служит разность между его точным значением $X_1 X_2$ и его приближённым значением $A_1 A_2$. Таким образом, мы можем написать:

$$\epsilon = \pm \epsilon_1 A_2 \pm \epsilon_2 A_1.$$

Переходя к предельной абсолютной ошибке $\epsilon_{\text{пр}}$, имеем:

$$\epsilon_{\text{пр}} = \pm (|\epsilon_1 A_2| + |\epsilon_2 A_1|).$$

Таким же приёмом можно найти выражения предельных абсолютных ошибок произведения трёх и в общем случае нескольких приближённых чисел. Но эти формулы оказываются достаточно сложными и при вычислении абсолютных ошибок произведения на практике не находят применения, так как в данном случае абсолютные ошибки можно вычислить значительно проще, пользуясь их общим выражением (формула (5'), стр. 27). В данном случае

$$\epsilon_{\text{пр}} = \delta_{\text{пр}} \cdot A,$$

где

$$A = A_1 A_2 A_3 \dots$$

Отсюда на основании формулы (15) находим:

$$\epsilon_{\text{пр}} = \pm (|\delta_1| + |\delta_2| + |\delta_3| + \dots) \cdot A_1 A_2 A_3 \dots \quad (16)$$

Эта формула даёт выражение предельной абсолютной ошибки произведения, выраженной через предельную относительную ошибку и произведение величин A .

Отсюда следует, что при определении ошибок произведения приближённых чисел необходимо сначала вычислить предельную относительную ошибку произведения (формула (15)), затем следует найти его предельную абсолютную ошибку (формула (16)).

Следует ещё указать величину тех ошибок, которые получаются в результате умножения приближённого числа на некоторое точное число N . С этой целью, ограничиваясь двумя множителями, находим по формулам (15) и (16):

$$\delta_{\text{пр}} = \pm (|\delta_1| + |\delta_2|),$$

$$\epsilon_{\text{пр}} = \pm A_1 A_2 (|\delta_1| + |\delta_2|) = \pm (|\epsilon_1 A_2| + |\epsilon_2 A_1|).$$

Если один из этих множителей, например A_2 , является точным числом, то его ошибки, т. е. величины δ_2 и ϵ_2 , равны нулю. Вследствие этого из последних формул, опуская в них индексы и заменяя A_2 через N , находим:

$$\delta_{\text{пр}} = \pm \delta,$$

$$\epsilon_{\text{пр}} = \pm N \cdot \epsilon.$$

Таким образом, мы вправе сказать, что:

При умножении приближённого числа на точное число N абсолютная ошибка произведения возрастает пропорционально N , а относительная ошибка произведения остаётся без изменения.

II. Ошибки частного. Два приближённых числа, данных в общем виде:

$$X_1 = A_1 \pm \epsilon_1,$$

$$X_2 = A_2 \pm \epsilon_2,$$

делим одно на другое. Имеем:

$$\frac{X_1}{X_2} = \frac{A_1 \pm \varepsilon_1}{A_2 \pm \varepsilon_2}. \quad (17)$$

Это выражение даёт возможность определить ошибки частного двух приближённых значений некоторых величин, причём сначала удобнее найти относительную ошибку частного.

1) Относительная ошибка частного. Переходя в формуле (17) к относительным ошибкам, имеем:

$$\frac{X_1}{X_2} = \frac{A_1}{A_2} \cdot \frac{1 \pm \delta_1}{1 \pm \delta_2}.$$

Отсюда, применяя соответствующую формулу для приближённых вычислений (стр. 373), находим:

$$\frac{X_1}{X_2} = \frac{A_1}{A_2} (1 \pm \delta_1 \mp \delta_2).$$

Из этого выражения мы видим, что относительная ошибка частного δ определяется величиной

$$\delta = \pm \delta_1 \mp \delta_2.$$

Отсюда находим предельную относительную ошибку частного:

$$\delta_{\text{пр}} = \pm (|\delta_1| + |\delta_2|). \quad (18)$$

Таким образом, мы можем сказать, что:

Предельная относительная ошибка частного равняется сумме абсолютных значений относительных ошибок делимого и делителя.

2) Абсолютная ошибка частного. Из формулы (17) можно определить и величину абсолютной ошибки частного. Для этого следует выполнить деление в правой части этой формулы. В результате этого деления мы получаем бесконечный ряд, в котором все члены, начиная с четвёртого, являются малыми величинами второго и более высоких порядков по отношению к частному $\frac{A_1}{A_2}$. Отбрасывая эти члены, находим после неболь-

шого преобразования, что абсолютная ошибка ε частного определяется выражением

$$\varepsilon = \frac{\pm \varepsilon_1 A_2 \mp \varepsilon_2 A_1}{A_2^2}.$$

Отсюда находим величину предельной абсолютной ошибки $\varepsilon_{\text{пр}}$ частного:

$$\varepsilon_{\text{пр}} = \pm \frac{|\varepsilon_1| A_2 + |\varepsilon_2| A_1}{A_2^2}. \quad (19)$$

Это выражение настолько сложно, что на практике им не пользуются и находят предельную абсолютную ошибку частного тем же приёмом, который был указан для нахождения предельной абсолютной ошибки произведения, т. е. предельную относительную ошибку, вычисленную по формуле (18), умножают на величину A_1/A_2 . Таким образом, можно написать:

$$\varepsilon_{\text{пр}} = \pm (|\delta_1| + |\delta_2|) \cdot \frac{A_1}{A_2}. \quad (20)$$

Подставляя в это выражение вместо δ_1 и δ_2 абсолютные ошибки ε_2 и ε_1 , находим:

$$\begin{aligned} \varepsilon_{\text{пр}} &= \pm \left(\left| \frac{\varepsilon_1}{A_1} \right| + \left| \frac{\varepsilon_2}{A_2} \right| \right) \frac{A_1}{A_2} = \pm \left(\left| \frac{\varepsilon_1}{A_2} \right| + \left| \frac{\varepsilon_2 A_1}{A_2^2} \right| \right) = \\ &= \pm \frac{|\varepsilon_1| A_2 + |\varepsilon_2| A_1}{A_2^2}, \end{aligned}$$

т. е. формулу (19).

Ошибки частного приближённых значений величин, как мы видим, следует находить в той же последовательности, как и ошибки произведения, т. е. сначала следует вычислить предельную относительную ошибку частного (формула (18)), затем следует найти его предельную абсолютную ошибку (формула (20)).

Следует ещё определить величину ошибок, которые получаются в результате деления приближённого числа на некоторое точное число N . Для этого мы можем воспользоваться формулами (18) и (19). Полагая в них

$$\begin{aligned} A_2 &= N, \\ \varepsilon_2 &= \delta_2 = 0, \end{aligned}$$

и опуская излишние индексы, находим:

$$\delta_{\text{пр}} = \pm \delta, \quad (18')$$

$$\epsilon_{\text{пр}} = \pm \frac{\epsilon}{N}. \quad (19')$$

Таким образом, мы видим, что:

При делении приближённого числа на точное число N абсолютная ошибка частного уменьшается в N раз, а его относительная ошибка остаётся без изменения.

5. Возведение в степень приближённых чисел и извлечение из них корня. I. Ошибки степени. Приближённое число, взятое в общем виде:

$$X = A \pm \epsilon,$$

возводим в некоторую степень n :

$$X^n = (A \pm \epsilon)^n. \quad (21)$$

Из этой формулы находим ошибки степени, причём удобнее найти сначала её относительную ошибку.

1) Относительная ошибка степени. Вводя в выражение (21) относительную ошибку, находим:

$$\begin{aligned} X^n &= A^n (1 \pm \delta)^n = \\ &= A^n \left(1 \pm n\delta + \frac{n(n-1)}{1 \cdot 2} \delta^2 \pm \frac{n(n-1)(n-2)}{1 \cdot 2 \cdot 3} \delta^3 + \dots \right), \end{aligned}$$

или с точностью до малых величин второго порядка:

$$X^n = A^n (1 \pm n\delta).$$

Отсюда мы видим, что относительная ошибка степени $\delta_{\text{ст}}$ определяется простым выражением

$$\delta_{\text{ст}} = \pm n\delta, \quad (22)$$

причём в данном случае нет оснований находить особое выражение для предельной ошибки. Таким образом приходим к следующему выводу:

Относительная ошибка степени равняется показателю степени, умноженному на относительную ошибку основания.

2) Абсолютная ошибка степени. Выражение абсолютной ошибки степени $\epsilon_{\text{ст}}$ можно получить из формулы (22). Подставляя в эту формулу вместо $\delta_{\text{ст}}$ и δ их выражения через абсолютные ошибки $\epsilon_{\text{ст}}$ и ϵ , находим:

$$\frac{\epsilon_{\text{ст}}}{A^n} = \pm n \frac{\epsilon}{A},$$

или

$$\epsilon_{\text{ст}} = \pm n A^{n-1} \epsilon. \quad (23)$$

При этом вновь следует иметь в виду, что находить особое выражение для предельной абсолютной ошибки в данном случае нет оснований. Таким образом, мы видим, что:

Абсолютная ошибка степени определяется произведением трёх множителей: показателя степени, основания, в степени на единицу ниже, и абсолютной ошибки основания.

Формулу (23) можно получить и непосредственно из выражения (21). Действительно, из этого выражения имеем:

$$\begin{aligned} X^n &= A^n \pm n A^{n-1} \epsilon + \frac{n(n-1)}{1 \cdot 2} A^{n-2} \epsilon^2 \pm \\ &\quad \pm \frac{n(n-1)(n-2)}{1 \cdot 2 \cdot 3} A^{n-3} \epsilon^3 + \dots \end{aligned}$$

Все члены этого ряда, начиная с третьего, являются малыми величинами второго и более высоких порядков по отношению к первому члену ряда A^n ; вследствие этого можем положить:

$$X^n = A^n \pm n A^{n-1} \epsilon.$$

Отсюда для абсолютной ошибки степени непосредственно получается выражение (23).

Следует ещё заметить, что формулы (22) и (23) можно получить также из формул соответственно (15) и (16), рассматривая возвышение в степень как умножение оди-

наковых по величине множителей. Но такой переход от формул (15) и (16) имеет смысл только для целых значений n , тогда как формулы (22) и (23) остаются справедливыми и для дробных показателей степени.

Формула (23) оказывается сложной для вычислений, поэтому для определения абсолютной ошибки степени на практике удобнее пользоваться приёмом, указанным выше, т. е. относительную ошибку степени, вычисленную по формуле (22), следует умножить на величину A^n ; таким образом, получаем:

$$\epsilon_{\text{ст}} = n\delta \cdot A^n. \quad (24)$$

Нетрудно видеть, что, подставляя в правой части этой формулы вместо δ её выражение через абсолютную ошибку ϵ , мы возвращаемся к формуле (23).

Ошибки степени следует находить в определённой последовательности: сначала следует найти по формуле (22) относительную ошибку степени, затем определить её абсолютную ошибку по формуле (24).

II. *Ошибки корня.* Возможность применять формулы (22) и (23) при дробных значениях n позволяет пользоваться этими формулами и при определении ошибок корня. Но вместе с тем правильность этих формул при определении ошибок корня нетрудно подтвердить и непосредственным выводом. Для этого следует из приближённого числа, взятого в общем виде, извлечь корень n -й степени:

$$\sqrt[n]{\bar{X}} = \sqrt[n]{A \pm \epsilon},$$

причём удобнее сначала определить относительную ошибку.

1) *Относительная ошибка корня.* Переходя в последней формуле к относительной ошибке, имеем:

$$\sqrt[n]{\bar{X}} = \sqrt[n]{A} \cdot \sqrt[n]{1 \pm \delta}.$$

Отсюда, применяя соответствующую формулу для приближённых вычислений, находим:

$$\sqrt[n]{\bar{X}} = \sqrt[n]{A} \left(1 \pm \frac{1}{n} \delta \right).$$

Эта формула показывает, что относительная ошибка δ_k корня определяется выражением

$$\delta_k = \pm \frac{1}{n} \delta, \quad (25)$$

причём выводить особое выражение для предельной относительной ошибки в данном случае нет оснований. Таким образом, мы видим, что:

Относительная ошибка корня из приближённого значения некоторой величины равна относительной ошибке этого числа, разделённой на степень корня.

2) *Абсолютная ошибка корня.* Подставляя в формулу (25) вместо δ_k и δ их выражения через абсолютные ошибки ϵ_k и ϵ , находим:

$$\frac{\epsilon_k}{\sqrt[n]{A}} = \pm \frac{1}{n} \cdot \frac{\epsilon}{A}.$$

Отсюда имеем:

$$\epsilon_k = \pm \frac{1}{n} A^{\frac{1}{n}-1} \cdot \epsilon, \quad (26)$$

причём в данном случае также нет оснований выводить особое выражение для предельной абсолютной ошибки. Следует заметить, что формулу (26) можно получить из уравнения (24), рассматривая извлечение корня как возведение в дробную степень, в данном случае равную $1/n$.

Формула (26) для вычислений является сложной, и при определении абсолютной ошибки корня на практике удобнее пользоваться другим выражением, которое мы получаем указанным выше приёмом, т. е. относительную ошибку корня, вычисленную по формуле (25), следует умножить на величину $\sqrt[n]{A}$. В результате этого получаем:

$$\epsilon_k = \pm \frac{1}{n} \cdot \sqrt[n]{A} \cdot \delta. \quad (27)$$

Отсюда следует, что ошибки корня необходимо определять в той же последовательности, как и ошибки степени: сначала вычислить по формуле (25) относительную ошибку корня, затем найти его абсолютную ошибку по формуле (27).

6. Несколько практических указаний. Часто можно слышать мнение, что математик, хорошо владеющий теорией, не может встретить больших трудностей в том, чтобы быстро и точно провести до конца любое вычисление. Но каждый, кто знаком с практикой вычислительных работ, знает, что они требуют громадной опытности и что на них приходится затрачивать несравненно больше внимания, чем это может показаться с первого взгляда. Действительно, получить при сложных вычислениях вполне безукоризненный результат удаётся только после затраты весьма большого труда, но вместе с тем всё же необходимо указать, что очень большая часть вины в этом часто падает почти исключительно на вычисляющего. Это объясняется главным образом тем, что вычисления очень часто выполняются без всякой системы, в полном беспорядке, не учитывается точность промежуточных результатов и лишние знаки не отбрасываются; вследствие этого цифры быстро нарастают, и вычисляющий начинает убеждаться в том, что он окончательно запутался. Не окончив вычислений, ему приходится их бросать с тем, чтобы через некоторое время попытаться вновь повторить тот же беспорядочный путь сначала.

Такой непроизводительной работы можно почти совершенно избежать, если весь вычислительный процесс производить в определённой последовательности, пользуясь рядом общих правил. Конечно, навык, привычка, находчивость играют при вычислительной работе очень большую роль, но вместе с тем при всех вычислениях рекомендуется иметь в виду следующие общие правила.

1. Точность окончательного результата вычислений, вообще говоря, не может быть больше точности исходных данных, т. е. точности измерений. Таким образом, относительная ошибка результата вычислений должна отвечать относительным ошибкам исходных данных; некоторые несущественные исключения из этого положения (стр. 61 и 75) не нарушают его общности.

Отсюда мы делаем ряд практических выводов:

а) Необходимо прежде всего установить точность исходных данных. Если исходные данные являются результатами измерений, то необходимо определить их отно-

сительные ошибки, иначе мы можем недооценить возможную точность результата, или, наоборот, будем вводить при вычислениях лишние знаки и тем создавать себе лишние осложнения.

б) На основании величины относительных ошибок исходных данных необходимо, ещё не приступая к вычислениям, приближённо установить относительную ошибку их окончательного результата. Для этого следует пользоваться в простейших случаях правилами, изложенными в этой главе, а в более сложных случаях — положениями теории ошибок (гл. III).

в) Если исходные данные имеют различную величину относительных ошибок, то необходимо принимать во внимание прежде всего ошибку наименее точного из исходных данных, так как его ошибка определяет в основном и ошибку результата вычислений.

г) Окончательный результат вычислений следует давать с тем числом значащих цифр, которое отвечает его ошибке.

2. Необходимо твёрдо помнить, что мы должны давать себе отчёт в точности каждого написанного числа. Вследствие этого необходимо определять относительные ошибки результатов всех промежуточных вычислений. В промежуточных результатах следует отбрасывать все лишние значащие цифры, оставляя на одну цифру больше того числа их, которое отвечало бы относительной ошибке окончательного результата вычислений. Эта лишняя цифра отбрасывается только в окончательном результате. Таким образом, все промежуточные вычисления следует вести, предполагая, что относительная ошибка результата несколько меньше той, которую мы установили.

3. Так как вычисления требуют много времени и большого напряжения, то следует пользоваться всеми средствами для их ускорения и облегчения, т. е. таблицами квадратов, кубов, корней и т. д., переходить к логарифмированию, применять формулы для приближённых вычислений и т. п. Точно так же следует пользоваться графическими методами вычислений и различными вычислительными приборами, счётными линейками, арифмометрами

и т. д. Необходимо при этом иметь в виду, что выбор вспомогательных средств при вычислениях должен определяться их точностью: принято считать, что логарифмическая линейка обычного типа и графические методы вычислений дают точность в среднем около 0,3%, при вычислениях с четырёхзначными таблицами логарифмов точность повышается до 0,03%, арифмометры и счётные машины могут дать точность, значительно более высокую.

4. Весь вычислительный процесс требует непрерывной проверки, т. е. должен идти под непрерывным контролем, который служит прекрасным средством своевременно обнаруживать ошибки, просчёты, описки, всегда возможные при вычислениях. Формы контроля могут быть очень различными, но весьма надёжным средством контроля являются обычные приёмы, которые рекомендуются в арифметике, т. е. проверка при помощи обратных действий. При этом необходимо каждое арифметическое действие, которое мы закончили, немедленно проверить при помощи обратного действия, и только после этого можно переходить к дальнейшим вычислениям. На примере личной практики, достаточно продолжительной, автор имел возможность убедиться в чрезвычайной действенности этой формы контроля. Вместе с тем необходимо иметь в виду, что почти неизбежной формой контроля вычислений следует считать их повторение два раза, желательно с применением различных вспомогательных средств при первом и втором вычислениях. Кроме того, очень рекомендуется производить вычисление одновременно двум лицам независимо друг от друга, сопоставляя результаты всех промежуточных вычислений.

5. Наконец, необходимо обращать большое внимание на внешнее оформление всего вычислительного процесса. Следует отказаться раз навсегда от широко распространённой привычки пользоваться при вычислениях случайными листочками бумаги, которые исписываются цифрами вдоль и поперёк. Основным требованием при вычислительных работах является наглядное расположение всех

вычислений на одном большом листе бумаги, на котором все отдельные этапы вычислений должны располагаться один за другим в определённой последовательности, так чтобы их можно было без труда находить. Все числа, над которыми выполняются отдельные арифметические действия, необходимо выписывать, даже если применяется счётная линейка или арифмометры, затем эти числа следует подчеркнуть, указать знак действия и выписать полученный результат. Вспомогательные и контрольные вычисления можно выполнять или на отдельной части того же общего листа, или, в случае недостатка места, на отдельном листе, но располагать эти вычисления необходимо в той последовательности, которая определяется ходом всего вычисления. Если приходится отбрасывать лишние цифры или округлять числа, а также вносить в них исправления, то это необходимо делать лёгким перечёркиванием лишних цифр, указывая сверху новые цифры так, чтобы прежние цифры, в случае необходимости, можно было без труда прочитать. Вообще все записи следует вести очень аккуратно, без грубых помарок и исправлений и следить за вертикальным расположением цифр одинаковых разрядов. Из последних соображений следует при вычислениях всегда пользоваться бумагой, разлинованной по клеткам. При повторении вычислений необходимо весь процесс воспроизвести в той же последовательности на новом листе бумаги, сопоставляя результаты отдельных операций: это даёт возможность очень легко установить все ошибки или недочёты, если они были допущены.

Выполнение этих советов может показаться очень утомительным, так как многое в них кажется излишним, слишком мелочным, несущественным. Но все эти правила выработаны практикой, и практика показывает, что если вычисления производить так, как здесь рекомендуется, то ошибки при вычислениях исключаются почти совершенно.

между физическими величинами, где элементарные формулы были бы почти совершенно непригодными.

Пример 1. При вертикальном падении небольшого шарика в вязкой жидкости его предельная скорость v определяется формулой

$$v = \frac{2}{9} \frac{gr^2(d_0 - d)}{\eta},$$

где g обозначает ускорение силы тяжести, r и d_0 обозначают радиус шарика и его плотность, а d и η — плотность жидкости и её коэффициент внутреннего трения.

Пример 2. Скорость распространения $v_{e\varphi}$ необыкновенного луча внутри одноосного кристалла по направлению, составляющему угол φ с направлением оптической оси кристалла определяется формулой

$$v_{e\varphi} = \frac{v_0 v_e}{\sqrt{v_0^2 \cos^2 \varphi + v_e^2 \sin^2 \varphi}},$$

где v_0 обозначает скорость луча обыкновенного, а v_e — скорость луча необыкновенного в направлении, перпендикулярном к оси кристалла.

Пример 3. Распределение энергии в спектре абсолютно чёрного тела в зависимости от длины волны λ и от температуры чёрного тела определяется известной формулой

$$E_{\lambda, T} = \frac{c_1}{\lambda^5} \cdot \frac{1}{\frac{c_2}{e^{\lambda T}} - 1},$$

где c_1 и c_2 являются некоторыми постоянными, величина которых определяется из опыта, а T обозначает абсолютную температуру чёрного тела.

Пример 4. Между постоянной Планка h , зарядом электрона e , его массой m и массой ядра водородного атома M теоретически была установлена определённая зависимость, которая выражается формулой

$$h = \sqrt[3]{\frac{2\pi^2 m e^4}{\left(1 + \frac{m}{M}\right) R_H}},$$

где R_H обозначает постоянную Ридберга для водорода.

ГЛАВА III

ОШИБКИ ПРИБЛИЖЁННЫХ ЗНАЧЕНИЙ ФУНКЦИЙ И ОБЩАЯ ТЕОРИЯ ОШИБОК

1. Основные задачи теории ошибок. Элементарные формулы, рассмотренные во второй главе, дают возможность определять величину ошибок, которые получаются в результате простейших арифметических действий над приближёнными числами. Все эти формулы являются весьма полезными; к ним приходится обращаться очень часто и не только при оценке ошибок результата уже законченных вычислений, но также при определении возможных ошибок предполагаемого вычисления, так как по величине ошибок данных нам приближённых чисел мы для любого арифметического действия над ними всегда сумеем определить величину ошибок заранее, ещё не приступая к вычислениям. Можно определить также и ту точность, с которой следует вести вычисления, чтобы ошибка результата не выходила за определённые пределы.

Но, несмотря на такие широкие области применения, всё же оказывается, что одних этих элементарных формул было бы совершенно недостаточно. Это объясняется тем, что различные физические величины, как известно, связаны между собой разнообразными и часто очень сложными математическими соотношениями. Вследствие этого, если мы определяем числовое значение какой-либо величины, то нам приходится выполнять не одно, а ряд последовательных арифметических действий над приближёнными числами, часто достаточно сложных. Можно привести сколько угодно примеров такой сложной зависимости

Пример 5. Траектория космических лучей под действием магнитного поля испытывает определённое искривление, что доказывает наличие в космических лучах электрически заряженных частичек; исследуя траекторию этих частичек в магнитном поле, можно найти их энергию E , величина которой определяется такой формулой:

$$E = c e \sqrt{\left(\frac{m_0 c}{e}\right)^2 + (Hr)^2},$$

где c , e и m обозначают соответственно скорость света в пустоте, заряд частички и её массу, а H и r обозначают напряжённость магнитного поля и радиус окружности, по которой движется частичка в магнитном поле.

Числовое значение каждой физической величины в подобных формулах мы можем определить, если будем знать числовые значения других величин в данной формуле. Эти значения мы можем получить, частью производя соответствующие измерения, частью воспользовавшись табличными данными. После этого нам придётся выполнить соответствующий ряд арифметических действий с приближёнными числами; вследствие этого необходимо принять во внимание ошибки исходных данных и определить отсюда ошибки результата вычислений, который мы получаем после большого числа различных арифметических действий. Применять в таких сложных условиях элементарные формулы было бы малопродуктивной работой, так как каждая из этих формул даёт величину ошибки только для одного определённого арифметического действия, иначе говоря, нам пришлось бы применять эти формулы много раз подряд в различной последовательности. Это явилось бы вычислительной операцией, весьма сложной и утомительной, а иногда и невозможной, например при наличии в формулах тригонометрических функций, как это имеет место во втором примере.

Отсюда следует, что необходимо установить общие приёмы и найти общие формулы для определения ошибок результата вычислений в общем случае при любой зависимости между физическими величинами, т. е. независимо от того, какие арифметические действия и в какой последовательности нам приходится производить при вы-

числениях. Для решения этих вопросов были применены методы дифференциального исчисления, что оказалось весьма удобным и плодотворным. В результате была разработана общая математическая теория ошибок, и большая часть её основных вопросов в настоящее время получила очень простое и полное разрешение. Относительно того общего метода, который применяется в математической теории ошибок, в дополнение к тому, что было сказано во введении (стр. 13), можно сказать ещё следующее.

Если мы считаем, что физическая величина, численное значение которой мы определяем, является функцией остальных физических величин в данной формуле, то, переходя к привычным символам, мы можем написать общее выражение:

$$y = f(x_1, x_2, x_3, \dots, x_n). \quad (1)$$

В этом выражении независимые переменные $x_1, x_2, x_3, \dots, x_n$ обозначают те величины, которые нам необходимо измерить для того, чтобы затем из результатов этих измерений определить числовое значение функции y . При измерении величин $x_1, x_2, x_3, \dots, x_n$ мы неизбежно допускаем некоторые ошибки, и каждая из этих ошибок оказывает определённое влияние на ошибку функции y . Абсолютные ошибки измерений всегда являются очень малыми по сравнению с измеряемыми величинами (стр. 23); вследствие этого в математической теории ошибок абсолютные ошибки условно рассматривают как бесконечно малые приращения величин $x_1, x_2, x_3, \dots, x_n$, взятые с двойным знаком, и обозначают соответственно этому символами $\pm dx_1, \pm dx_2, \pm dx_3, \dots, \pm dx_n$. Абсолютная ошибка функции y , которая вызывается этими бесконечно малыми приращениями её аргументов, по основному свойству непрерывных функций является также бесконечно малой величиной, которую мы должны, очевидно, обозначить символом $\pm dy$. Таким образом, мы можем написать:

$$y \pm dy = f(x_1 \pm dx_1, x_2 \pm dx_2, x_3 \pm dx_3, \dots, x_n \pm dx_n). \quad (2)$$

По отношению к этому уравнению следует ещё раз напомнить, что применять его можно только к функциям

непрерывным. Таким образом, необходимо всегда предварительно убедиться в том, что функция f в уравнении (1) удовлетворяет условиям непрерывности.

Этот общий приём теории ошибок не свободен от упрека в одном отношении: ошибки измерений, которые мы принимаем за бесконечно малые величины, основному свойству этих величин не отвечают. Действительно, бесконечно малая величина является величиной переменной, имеющей своим пределом нуль. Ошибки измерений являются величинами малыми, но не переменными, а постоянными. Этот упрек, однако, можно считать несерьёзным хотя бы потому, что постоянную величину можно рассматривать как частное значение переменной величины. Вследствие этого мы будем считать, что применение уравнений (1) и (2) в теории ошибок является достаточно обоснованным.

В общей теории ошибок рассматриваются три основные задачи, которые можно охарактеризовать следующим образом:

1) Если мы определяем некоторую величину y для этого производим измерение нескольких величин $x_1, x_2, x_3, \dots, x_n$, то прежде всего мы должны иметь возможность находить ошибки y по величине ошибок, допущенных при измерениях $x_1, x_2, x_3, \dots, x_n$. Таким образом, эту первую задачу, которую часто называют прямой задачей теории ошибок, можно определить так: вычисление ошибок функции, если известны ошибки её аргументов и вид функциональной зависимости.

2) Если мы, определяя некоторую величину y , зависящую от нескольких других величин $x_1, x_2, x_3, \dots, x_n$, ставим условие, чтобы ошибки y не выходили за определённые пределы, то мы, очевидно, должны иметь возможность находить величину ошибок, которые можно допустить при измерении $x_1, x_2, x_3, \dots, x_n$, не нарушая поставленного условия. Таким образом, эту вторую задачу, которую часто называют обратной задачей теории ошибок, можно определить так: вычисление ошибок аргументов, если известны ошибки функции и вид функциональной зависимости.

3) Наконец, мы можем поставить ещё вопрос о нахождении лучших условий для измерений, т. е. если мы определяем некоторую величину y и для этого должны измерить величины $x_1, x_2, x_3, \dots, x_n$, то нельзя ли найти такие условия этих измерений, при которых определение y будет выполнено наиболее точно. Таким образом, эту третью задачу теории ошибок можно определить так: нахождение наиболее выгодных условий измерения, т. е. таких условий, при которых ошибка функции оказывается наименьшей.

Исследования показывают, что:

1) Прямая задача теории ошибок, т. е. определение ошибок функции в зависимости от ошибок её аргументов, во всех случаях находит вполне удовлетворительное решение.

2) Обратная задача теории ошибок, т. е. определение ошибок аргументов в зависимости от ошибок функции, оказывается более сложной, и её решение часто имеет несколько неопределённый характер.

3) Третья задача теории ошибок, т. е. нахождение наиболее выгодных условий измерения, далеко не во всех случаях допускает определённое решение. Основную роль играет здесь вид функции f , и для очень многих функций определённое решение оказывается принципиально невозможным.

Необходимо ещё заметить, что возможности теории ошибок не ограничиваются этими тремя основными задачами и что эту теорию можно применять при решении некоторых других вопросов, в частности при оценке сравнительных достоинств различных измерительных методов, как это мы увидим в дальнейшем.

2. Ошибки функции одного и двух независимых переменных. I. Одно независимое переменное.

В случае одного независимого переменного уравнения (1) и (2) получают вид:

$$y = f(x), \quad (1')$$

$$y \pm dy = f(x \pm dx). \quad (2')$$

Применяя к правой части уравнения (2') формулу Тейлора, имеем:

$$y \pm dy = f(x) \pm dx f'(x) + \frac{(dx)^2}{2} f''(x) \pm \frac{(dx)^3}{2 \cdot 3} f'''(x) + \dots$$

В правой части этого уравнения отбрасываем, по общему правилу приближённых вычислений (стр. 51), бесконечно малые величины второго и более высоких порядков.

Находим:

$$y \pm dy = f(x) \pm dx f'(x).$$

Отсюда, принимая во внимание уравнение (1'), получаем:

$$dy = \pm dx f'(x).$$

Из этого уравнения, обозначая абсолютные ошибки функции y и независимого переменного x соответственно через ϵ_y и ϵ_x , имеем:

$$\epsilon_y = \pm \epsilon_x f'(x). \quad (3)$$

Таким образом, мы видим, что:

Абсолютная ошибка функции одного независимого переменного определяется произведением абсолютной ошибки аргумента на производную этой функции.

Из уравнения (3) нетрудно определить величину относительной ошибки y . Действительно, разделив обе части этого уравнения на y , находим:

$$\frac{\epsilon_y}{y} = \pm \frac{\epsilon_x f'(x)}{y},$$

или

$$\frac{dy}{y} = \pm \frac{dx \frac{dy}{dx}}{f(x)} = \pm \frac{df(x)}{f(x)}.$$

В этом уравнении левая часть является относительной ошибкой y , которую обозначим через δ_y , а правая часть представляет собой дифференциал натурального

логарифма $f(x)$. Таким образом, можно написать:

$$\delta_y = \pm d \ln f(x). \quad (4)$$

Отсюда следует, что:

Относительная ошибка функции одного независимого переменного равняется дифференциалу натурального логарифма этой функции.

Из уравнений (3) и (4) мы видим, что для любой функции одного независимого переменного обе её ошибки можно находить при помощи очень простых приёмов, причём особенно простое выражение получается для относительной ошибки.

II. *Два независимых переменных.* В случае двух независимых переменных уравнения (1) и (2) получают вид:

$$y = f(x_1, x_2), \quad (1'')$$

$$y \pm dy = f(x_1 \pm dx_1, x_2 \pm dx_2). \quad (2'')$$

Вновь применяя к правой части уравнения (2'') формулу Тейлора, имеем:

$$\begin{aligned} y \pm dy = f(x_1, x_2) \pm & \left[\frac{\partial f(x_1, x_2)}{\partial x_1} dx_1 \pm \frac{\partial f(x_1, x_2)}{\partial x_2} dx_2 \right] + \\ & + \frac{1}{2} \left[\frac{\partial^2 f(x_1, x_2)}{\partial x_1^2} (dx_1)^2 \pm 2 \frac{\partial^2 f(x_1, x_2)}{\partial x_1 \partial x_2} dx_1 dx_2 + \right. \\ & \left. + \frac{\partial^2 f(x_1, x_2)}{\partial x_2^2} (dx_2)^2 \right] \pm \dots \end{aligned}$$

Отбрасывая в этом уравнении бесконечно малые величины второго и более высоких порядков и принимая во внимание уравнение (1''), находим:

$$dy = \pm \frac{\partial f(x_1, x_2)}{\partial x_1} dx_1 \pm \frac{\partial f(x_1, x_2)}{\partial x_2} dx_2.$$

В этом уравнении в левой части стоит абсолютная ошибка функции y , а правая часть представляет собой полный дифференциал этой функции, причём каждый из частных дифференциалов определяет величину той

части ошибки y , которую вносят ошибки dx_1 и dx_2 в отдельности. Так как обе эти частные ошибки имеют двойной знак, то мы переходим, как это делали и раньше, к вычислению предельной абсолютной ошибки y , предполагая, что имеет место наименее благоприятный случай, т. е. что обе частные ошибки имеют одинаковый знак. Таким образом, обозначая предельную абсолютную ошибку y через $(\epsilon_y)_{\text{пр}}$, можем написать:

$$(\epsilon_y)_{\text{пр}} = \pm \left\{ \left| \frac{\partial f(x_1, x_2)}{\partial x_1} dx_1 \right| + \left| \frac{\partial f(x_1, x_2)}{\partial x_2} dx_2 \right| \right\} \dots \quad (5)$$

Отсюда следует, что:

Предельная абсолютная ошибка функции двух независимых переменных определяется суммой абсолютных величин частных дифференциалов этой функции.

Для того чтобы найти предельную относительную ошибку y , делим левую часть уравнения (5) на y , а его правую часть на $f(x_1, x_2)$. В результате в левой части мы получаем предельную относительную ошибку y , которую обозначим через $(\delta_y)_{\text{пр}}$, а правая часть после небольших преобразований оказывается равной дифференциалу натурального логарифма $f(x_1, x_2)$. Таким образом, можно написать:

$$(\delta_y)_{\text{пр}} = \pm d \ln f(x_1, x_2). \quad (6)$$

Отсюда следует, что:

Предельная относительная ошибка функции двух независимых переменных равняется дифференциалу натурального логарифма этой функции, но при этом следует брать сумму абсолютных величин всех членов этого выражения.

Для того чтобы ознакомиться с применением формул (3), (4), (5) и (6), рассмотрим несколько примеров:

Пример 1. Для определения объема шара V была применена обычная формула:

$$V = \frac{4}{3} \pi r^3, \quad (7)$$

причем из измерений было найдено, что радиус шара r можно считать равным:

$$r = |5,012 \pm 0,005| \text{ см.}$$

Найдём объем шара V и величину его абсолютной и относительной ошибок ϵ_V и δ_V .

В данном случае ошибки радиуса r , абсолютная ϵ_r и относительная δ_r равны:

$$\epsilon_r = \pm 0,005 \text{ см и } \delta_r = \pm \frac{0,005}{5,012} = \pm 0,001 = \pm 0,1\%.$$

Абсолютную ошибку объема шара ϵ_V находим, применяя формулу (3):

$$\epsilon_V = \pm \epsilon_r \cdot \frac{4}{3} \pi \cdot 3r^2 = \pm \epsilon_r \cdot 4\pi r^2.$$

Подставляя числовые данные, имеем:

$$\epsilon_V = \pm 0,005 \cdot 4\pi \cdot 25,12 \text{ см}^3 = \pm 1,58 \text{ см}^3,$$

или, с округлением:

$$\epsilon_V = \pm 2 \text{ см}^3.$$

Вычисляя приближённое значение объема шара V , находим по формуле (7):

$$V = \frac{4}{3} \pi (5,012)^3 \text{ см}^3 = 527,37 \text{ см}^3.$$

Окончательный результат для V , принимая во внимание величину абсолютной ошибки ϵ_V , следует написать в таком виде:

$$V = (527 \pm 2) \text{ см}^3.$$

Относительную ошибку объема шара δ_V находим, применяя формулу (4):

$$\begin{aligned} \delta_V &= \pm d (\ln 4 - \ln 3 + \ln \pi + 3 \ln r) = \pm d (3 \ln r) = \\ &= \pm 3 \cdot \frac{dr}{r} = \pm 3\delta_r, \end{aligned}$$

или:

$$\delta_V = \pm 3 \cdot 0,001 = \pm 0,003 = \pm 0,3\%.$$

Таким образом, мы видим, что относительная ошибка V увеличилась в три раза по сравнению с относительной ошибкой r . Это объясняется тем, что в формуле (7) r входит в третьей степени, и вследствие этого относительная ошибка V должна утроиться, как то наглядно видно и из элементарной формулы (22) (стр. 72).

Пример 2. В одном из старинных типов гальванометров, который при некоторых измерениях применяется и в настоящее время, сила тока I и угол отклонения φ магнитной стрелки связаны соотношением

$$I = k \operatorname{tg} \varphi, \quad (8)$$

где k обозначает так называемую постоянную прибора, которая определяется его конструктивными данными.

При измерениях с одним из таких приборов были получены в двух случаях значения угла φ , равные:

$$\varphi' = 6^\circ 17' \pm 1'$$

и

$$\varphi'' = 43^\circ 32' \pm 1',$$

причём постоянная данного прибора, измеренная весьма точно, оказалась равной

$$k = 5,031 \cdot 10^{-7} a.$$

Определим для обоих случаев приближённые значения силы тока I' и I'' и их ошибки, абсолютные, ϵ'_I и ϵ''_I , и относительные, δ'_I и δ''_I .

Определяя абсолютные и относительные ошибки углов φ' и φ'' , находим:

$$\epsilon'_\varphi = \epsilon''_\varphi = \pm 1' \text{ или в радианах } \pm 0,0003,$$

$$\delta'_\varphi = \pm \frac{1'}{6^\circ 17'} = \pm 0,0027 = \pm 0,27\% \text{ и}$$

$$\delta''_\varphi = \pm \frac{1'}{43^\circ 32'} = \pm 0,0004 = \pm 0,04\%.$$

Находим абсолютные ошибки I' и I'' .

Общее выражение, определяющее абсолютную ошибку ϵ_I , в данном случае получаем из формулы (3):

$$\epsilon_I = \pm \epsilon_\varphi \frac{d(k \operatorname{tg} \varphi)}{d\varphi} = \pm \epsilon_\varphi \cdot k \frac{1}{\cos^2 \varphi}.$$

Соответственно этой формуле имеем:

$$\begin{aligned} \epsilon'_I &= \pm 0,0003 \cdot 5,031 \cdot 10^{-7} \cdot \frac{1}{\cos^2 6^\circ 17'} a = \\ &= \pm 0,0015 \cdot 10^{-7} a = \pm 0,001(5) \cdot 10^{-7} a, \end{aligned}$$

$$\begin{aligned} \epsilon''_I &= \pm 0,0003 \cdot 5,031 \cdot 10^{-7} \cdot \frac{1}{\cos^2 43^\circ 32'} a = \\ &= \pm 0,0029 \cdot 10^{-7} a = \pm 0,003 \cdot 10^{-7} a. \end{aligned}$$

Определяя по формуле (8) приближённое значение силы тока в первом и втором случаях, имеем:

$$I' = 5,031 \cdot 10^{-7} \operatorname{tg} 6^\circ 17' a = 0,5539 \cdot 10^{-7} a,$$

$$I'' = 5,031 \cdot 10^{-7} \operatorname{tg} 43^\circ 32' a = 4,4799 \cdot 10^{-7} a.$$

Принимая во внимание величину абсолютных ошибок ϵ'_I и ϵ''_I , мы видим, что окончательное значение I' и I'' следует написать в таком виде:

$$I' = (0,554 \pm 0,001(5)) \cdot 10^{-7} a,$$

$$I'' = (4,480 \pm 0,003) \cdot 10^{-7} a.$$

Находим относительные ошибки I' и I'' . Общее выражение относительной ошибки в данном случае имеет вид:

$$\begin{aligned} \delta_I &= \pm d \ln k \operatorname{tg} \varphi = \pm d \ln \operatorname{tg} \varphi = \\ &= \pm \frac{d \operatorname{tg} \varphi}{\operatorname{tg} \varphi} = \pm \frac{2d\varphi}{\sin 2\varphi} = \pm \frac{2\epsilon_\varphi}{\sin 2\varphi}. \end{aligned}$$

Соответственно этому имеем:

$$\delta'_I = \pm \frac{2 \cdot 0,0003}{\sin 12^\circ 34'} = \pm 0,003 = \pm 0,3\%,$$

$$\delta''_I = \pm \frac{2 \cdot 0,0003}{\sin 87^\circ 41'} = \pm 0,0006 = \pm 0,06\%.$$

Следует обратить внимание на то, что при функциональной зависимости между двумя величинами, которая представляется формулой (8), величина ошибок функции I , в данном случае силы тока, зависит от значения аргумента φ , т. е. от угла отклонения магнитной стрелки. Но эта зависимость для абсолютных и для относительных ошибок оказывается

различной: действительно, при изменении φ в пределах приближённо от 6 до 45° абсолютная ошибка силы тока I возрастает (в два раза), а её относительная ошибка уменьшается (в пять раз).

Пример 3. Формула диффракционной решётки имеет такой вид:

$$\lambda = \frac{1}{n} D \sin \varphi, \quad (9)$$

где λ и φ обозначают длину волны светового луча и угол его отклонения от первоначального направления, а n и D обозначают порядок диффракционного спектра и так называемую постоянную решётки, т. е. расстояние между серединами двух соседних штрихов.

При измерении длины волны светового луча применялась решётка, постоянная которой D , очень точно измеренная, оказалась равной

$$D = 3,996 \cdot 10^{-4} \text{ см.}$$

Измерение угла φ отклонения луча в спектре второго порядка, т. е. для n , равного двум, дало такие результаты:

$$\varphi = 17^\circ 9' \pm 1'.$$

Найдём длину волны светового луча и её ошибки ϵ_λ и δ_λ . Абсолютная и относительная ошибки угла φ , очевидно, равны:

$$\epsilon_\varphi = \pm 1' \text{ или в радианах } \pm 0,0003, \\ \delta_\varphi = \frac{1'}{17^\circ 9'} = 0,001 = 0,1\%.$$

Определяем абсолютную ошибку длины волны ϵ_λ .

Общее выражение абсолютной ошибки в данном случае имеет вид

$$\epsilon_\lambda = \epsilon_\varphi \cdot \frac{d\left(\frac{1}{n} \cdot D \sin \varphi\right)}{d\varphi} = \epsilon_\varphi \cdot \frac{1}{n} \cdot D \cdot \cos \varphi.$$

Подставляя числовые данные, находим:

$$\epsilon_\lambda = \pm 0,0003 \cdot \frac{1}{2} \cdot 3,996 \cdot 10^{-4} \cos 17^\circ 9' \text{ см} = \\ = \pm 5,7 \cdot 10^{-8} \text{ см} = \pm 0,0006 \mu.$$

По формуле (9) находим длину волны λ светового луча:

$$\lambda = \frac{1}{2} \cdot 3,996 \cdot 10^{-4} \sin 17^\circ 9' \text{ см} = 0,58921 \mu.$$

Этот результат в соответствии с величиной абсолютной ошибки ϵ_λ следует написать в таком виде:

$$\lambda = (0,5892 \pm 0,0006) \mu.$$

Находим относительную ошибку длины волны δ_λ .

Общее выражение относительной ошибки в данном случае имеет вид

$$\delta_\lambda = \pm d \ln \frac{1}{n} \cdot D \cdot \sin \varphi = \pm d \ln \sin \varphi = \\ = \pm \frac{\cos \varphi}{\sin \varphi} d\varphi = \pm \operatorname{ctg} \varphi \cdot \epsilon_\varphi.$$

Подставляя числовые данные, находим:

$$\delta_\lambda = \pm \operatorname{ctg} 17^\circ 9' \cdot 0,0003 = \pm 0,00097 = \pm 0,001 = \pm 0,1\%.$$

В данном случае относительные ошибки результата вычислений (δ_λ) и исходных (δ_φ) оказываются одинаковыми.

Пример 4. При определении ускорения силы тяжести g методом обратного маятника применяется формула:

$$g = \pi^2 \frac{l}{t^2}, \quad (10)$$

где l и t обозначают приведённую длину маятника и период его простого колебания.

Измерение величин l и t при одном из таких определений g дало следующие результаты:

$$l = (50,02 \pm 0,01) \text{ см}$$

и

$$t = (0,7098 \pm 0,0001) \text{ сек.}$$

Найдём по этим данным ускорение силы тяжести g и его предельные ошибки, абсолютную (ϵ_g) и относительную (δ_g).

Абсолютные и относительные ошибки l и t , очевидно, равны:

$$\begin{aligned}\epsilon_l &= \pm 0,01 \text{ см}, & \delta_l &= \pm \frac{0,01}{50,02} = \pm 0,0002 = \pm 0,02\%, \\ \epsilon_t &= \pm 0,0001 \text{ сек}, & \delta_t &= \pm \frac{0,0001}{0,7098} = \pm 0,00014 = \pm 0,014\%.\end{aligned}$$

Так как мы имеем функцию двух независимых переменных, то следует применить формулы (5) и (6).

Для определения абсолютной ошибки g находим согласно формуле (5) частные дифференциалы функции (10) по l и t :

$$\begin{aligned}\partial g_l &= \frac{\partial \left(\pi^2 \frac{l}{t^2} \right)}{\partial l} \cdot dl = \frac{\pi^2}{t^2} dl, \\ \partial g_t &= \frac{\partial \left(\pi^2 \frac{l}{t^2} \right)}{\partial t} \cdot dt = -2 \cdot \frac{\pi^2 l}{t^3} dt.\end{aligned}$$

Отсюда, применяя формулу (5), получаем:

$$(\epsilon_g)_{\text{пр}} = \pm \left(\frac{\pi^2}{t^2} \cdot dl + 2 \cdot \frac{\pi^2 l}{t^3} dt \right).$$

При подстановке числовых данных следует определять размерность отдельных членов, так как анализ размерностей является здесь лучшим методом для проверки произведённых алгебраических преобразований. Вследствие этого находим:

$$\begin{aligned}(\epsilon_g)_{\text{пр}} &= \pm \left(\frac{\pi^2}{(0,7098)^2} \cdot 0,01 \frac{\text{см}}{\text{сек}^2} + 2 \cdot \frac{\pi^2 \cdot 50,02}{(0,7098)^3} \cdot 0,0001 \frac{\text{см}}{\text{сек}^2} \right) = \\ &= \pm \left(0,197 \frac{\text{см}}{\text{сек}^2} + 0,277 \frac{\text{см}}{\text{сек}^2} \right) = \pm 0,47 \frac{\text{см}}{\text{сек}^2} = \pm 0,5 \frac{\text{см}}{\text{сек}^2}.\end{aligned}$$

Величину g определяем из формулы (10):

$$g = \pi^2 \cdot \frac{50,02}{(0,7098)^2} \frac{\text{см}}{\text{сек}^2} = 979,87 \frac{\text{см}}{\text{сек}^2}.$$

Этот результат, принимая во внимание величину абсолютной ошибки ϵ_g , следует написать в таком виде:

$$g = (979,9 \pm 0,5) \frac{\text{см}}{\text{сек}^2}.$$

Находим предельную относительную ошибку ускорения силы тяжести $(\delta_g)_{\text{пр}}$. Из формулы (6) имеем:

$$\begin{aligned}(\delta_g)_{\text{пр}} &= \pm d \ln \pi^2 \frac{l}{t^2} = \pm (d \ln l + 2d \ln t) = \\ &= \pm \left(\frac{dl}{l} + 2 \frac{dt}{t} \right) = \pm (\delta_l + 2\delta_t).\end{aligned}$$

Переходя к числовым данным, находим:

$$(\delta_g)_{\text{пр}} = \pm (0,02 + 2 \cdot 0,014)\% = \pm 0,048\% = \pm 0,05\%.$$

Следует обратить внимание на то обстоятельство, что хотя относительная ошибка t меньше относительной ошибки l , однако её влияние на ошибку g оказывается больше. Это объясняется тем, что в формуле (10) период t входит во второй степени, и его ошибка в окончательном результате удваивается.

Пример 5. Формула дисперсионной трёхгранной призмы, установленной на угол наименьшего отклонения, имеет вид:

$$n = \frac{\sin \frac{1}{2} (\alpha + \beta)}{\sin \frac{1}{2} \alpha}, \quad (11)$$

где n , α и β обозначают соответственно коэффициент преломления, преломляющий угол призмы и угол наименьшего отклонения.

При определении коэффициента преломления n на спектрометре с такой призмой, изготовленной из оптического стекла, измерялись углы α и β . Их значения оказались:

$$\alpha = 60^\circ 12,5' \pm 0,5',$$

$$\beta = 41^\circ 26,25' \pm 0,5'.$$

Найдём значение коэффициента преломления n и его предельные ошибки, абсолютную $(\epsilon_n)_{\text{пр}}$ и относительную $(\delta_n)_{\text{пр}}$.

Ошибки углов α и β , абсолютные, ϵ_α и ϵ_β , и относительные, δ_α и δ_β , очевидно, равны:

$$\epsilon_\alpha = \epsilon_\beta = \pm 0,5' \text{ или в радианах } \pm 0,000145,$$

$$\delta_\alpha = \pm \frac{0,5'}{60^\circ 12,5'} = \pm 0,00014 = \pm 0,014\%,$$

$$\delta_\beta = \pm \frac{0,5'}{41^\circ 26,25'} = \pm 0,0002 = \pm 0,02\%.$$

При определении $(\epsilon_n)_{\text{пр}}$ функцию (11) удобнее представить в таком виде:

$$n = \cos \frac{1}{2} \beta + \operatorname{ctg} \frac{1}{2} \alpha \sin \frac{1}{2} \beta. \quad (11')$$

Находим значение углов $\frac{1}{2} \alpha$ и $\frac{1}{2} \beta$, а также их ошибки; кроме того, находим значение угла $\frac{1}{2} (\alpha + \beta)$ и его абсолютную ошибку, что нам в дальнейшем окажется необходимым. При определении ошибок этих углов необходимо иметь в виду формулы (9) (стр. 59) и (18') и (19') (стр. 72).

Получаем:

$$\frac{1}{2} \alpha = 30^\circ 6,25', \quad \frac{1}{2} \beta = 20^\circ 43,12', \quad \frac{1}{2} (\alpha + \beta) = 50^\circ 49,37',$$

$$\epsilon_{\frac{\alpha}{2}} = \epsilon_{\frac{\beta}{2}} = \pm 0,25' \text{ или в радианах } \pm 0,000072,$$

$$\epsilon_{\frac{(\alpha+\beta)}{2}} = \pm 0,5' \text{ или в радианах } \pm 0,000145,$$

$$\delta_{\frac{\alpha}{2}} = \pm 0,00014 = \pm 0,014\%,$$

$$\delta_{\frac{\beta}{2}} = \pm 0,0002 = \pm 0,02\%.$$

Так как n является функцией двух независимых переменных, то при вычислении $(\epsilon_n)_{\text{пр}}$ применяем формулу (5). Определяя из уравнения (11') значение частных

дифференциалов n по $\frac{1}{2} \alpha$ и по $\frac{1}{2} \beta$, имеем:

$$\partial n_{\frac{\alpha}{2}} = - \frac{\sin \frac{1}{2} \beta}{\sin^2 \frac{1}{2} \alpha} \cdot d \left(\frac{1}{2} \alpha \right) = - \frac{\sin \frac{1}{2} \beta}{\sin^2 \frac{1}{2} \alpha} \cdot \epsilon_{\frac{\alpha}{2}},$$

$$\begin{aligned} \partial n_{\frac{\beta}{2}} &= - \sin \frac{1}{2} \beta \cdot d \left(\frac{1}{2} \beta \right) + \operatorname{ctg} \frac{1}{2} \alpha \cdot \cos \frac{1}{2} \beta \cdot d \left(\frac{1}{2} \beta \right) = \\ &= \left(- \sin \frac{1}{2} \beta + \operatorname{ctg} \frac{1}{2} \alpha \cdot \cos \frac{1}{2} \beta \right) \cdot \epsilon_{\frac{\beta}{2}}. \end{aligned}$$

Отсюда находим предельную абсолютную ошибку n :

$$(\epsilon_n)_{\text{пр}} = \pm \left[\frac{\sin \frac{1}{2} \beta}{\sin^2 \frac{1}{2} \alpha} \cdot \epsilon_{\frac{\alpha}{2}} + \left(\sin \frac{1}{2} \beta + \operatorname{ctg} \frac{\alpha}{2} \cdot \cos \frac{1}{2} \beta \right) \cdot \epsilon_{\frac{\beta}{2}} \right].$$

Переходя к числовым данным, получаем:

$$\begin{aligned} (\epsilon_n)_{\text{пр}} &= \pm \left(\frac{\sin 20^\circ 43,12'}{\sin^2 30^\circ 6,25'} + \sin 20^\circ 43,12' + \right. \\ &\quad \left. + \operatorname{ctg} 30^\circ 6,25' \cdot \cos 20^\circ 43,12' \right) 0,000072. \end{aligned}$$

Производя вычисления, находим:

$$(\epsilon_n)_{\text{пр}} = \pm 0,00026 = \pm 0,0003.$$

Из формулы (11) находим n :

$$n = \frac{\sin \frac{1}{2} (60^\circ 12,5' + 41^\circ 26,25')}{\sin \frac{1}{2} 60^\circ 12,5'} = 1,54556.$$

Этот результат, принимая во внимание значение $(\epsilon_n)_{\text{пр}}$, следует написать в таком виде:

$$n = 1,5456 \pm 0,0003.$$

При определении $(\delta_n)_{\text{пр}}$ применяем формулу (6). Имеем:

$$\begin{aligned} (\delta_n)_{\text{пр}} &= \pm d \ln \frac{\sin \frac{1}{2} (\alpha + \beta)}{\sin \frac{1}{2} \alpha} = \\ &= \pm d \left(\ln \sin \frac{1}{2} (\alpha + \beta) + \ln \sin \frac{1}{2} \alpha \right). \end{aligned}$$

Выполняя дифференцирование, находим:

$$\begin{aligned} (\delta_n)_{\text{пр}} &= \pm \left(\operatorname{ctg} \frac{1}{2} (\alpha + \beta) \cdot d \frac{1}{2} (\alpha + \beta) + \operatorname{ctg} \frac{1}{2} \alpha \cdot d \frac{1}{2} \alpha \right) = \\ &= \pm \left(\operatorname{ctg} \frac{1}{2} (\alpha + \beta) \cdot \varepsilon_{\frac{\alpha + \beta}{2}} + \operatorname{ctg} \frac{1}{2} \alpha \cdot \varepsilon_{\frac{\alpha}{2}} \right). \end{aligned}$$

Подставляя числовые данные, получаем:

$$\begin{aligned} (\delta_n)_{\text{пр}} &= \pm (\operatorname{ctg} 50^\circ 49,37' \cdot 0,000145 + \operatorname{ctg} 30^\circ 6,25' \cdot 0,000072) = \\ &= \pm 0,00024 = \pm 0,02\%. \end{aligned}$$

Во всех этих примерах обращает на себя внимание то обстоятельство, что определение абсолютных ошибок функций является более сложной операцией, так как при определении относительных ошибок функций все вычисления, как правило, оказываются более простыми. Кроме того, из этих примеров видно, какое влияние на ошибку результата оказывает вид функциональной зависимости f . В частности, несколько своеобразные выражения ошибок получаются в тех случаях, когда аргумент входит под знаком тригонометрической функции, как, например, в примерах 2, 3 и 5. Вследствие этого очень полезно заранее найти выражения абсолютных и относительных ошибок для определённых функций; некоторые из таких формул даны в приложении (стр. 374).

3. Ошибки функции нескольких независимых переменных. Случай функции двух независимых переменных, который представлен формулами (5) и (6), можно распространить и на общий случай функции нескольких независимых переменных. Действительно, основное уравнение в этом случае имеет вид

$$y = f(x_1, x_2, x_3, \dots, x_n).$$

Рассуждая совершенно так же, как и в случае двух независимых переменных, мы приходим к выводу, что предельные ошибки функции y в общем случае, абсолютная $(\varepsilon_y)_{\text{пр}}$ и относительная $(\delta_y)_{\text{пр}}$, должны определяться соответственно выражениями

$$\begin{aligned} (\varepsilon_y)_{\text{пр}} &= \pm \left[\left| \frac{\partial f(x_1, x_2, x_3, \dots, x_n)}{\partial x_1} dx_1 \right| + \right. \\ &\quad \left. + \left| \frac{\partial f(x_1, x_2, x_3, \dots, x_n)}{\partial x_2} dx_2 \right| + \left| \frac{\partial f(x_1, x_2, x_3, \dots, x_n)}{\partial x_3} dx_3 \right| + \dots \right. \\ &\quad \left. \dots + \left| \frac{\partial f(x_1, x_2, x_3, \dots, x_n)}{\partial x_n} dx_n \right| \right], \quad (5') \end{aligned}$$

$$(\delta_y)_{\text{пр}} = \pm d [\ln f(x_1, x_2, x_3, \dots, x_n)]. \quad (6')$$

Таким образом, мы можем сказать, что:

1) Предельная абсолютная ошибка функции нескольких независимых переменных определяется суммой абсолютных величин всех частных дифференциалов этой функции и

2) Предельная относительная ошибка функции нескольких независимых переменных равняется дифференциалу натурального логарифма этой функции, причём следует брать сумму абсолютных величин всех членов этого выражения.

Для того чтобы ознакомиться с применением этих формул на практике, рассмотрим несколько примеров.

Пример 1. При определении удельной теплоёмкости c твёрдых тел методом водяного калориметра часто пользуются элементарной формулой:

$$c = \frac{M(\theta - t_0)}{m(t_1 - \theta)}, \quad (12)$$

где M и m обозначают массу воды в калориметре и массу исследуемого тела, а t_0 , t_1 и θ обозначают соответственно начальную температуру воды в калориметре, температуру тела m до введения его в калориметр и окончательную температуру воды в калориметре.

При определении этим методом удельной теплоёмкости алюминия были получены такие значения величин M , m , t_0 , t_1 и θ :

$$M = (250,0 \pm 0,2) \text{ г},$$

$$m = (62,31 \pm 0,02) \text{ г},$$

$$t_0 = 13,52^\circ \pm 0,01^\circ,$$

$$t_1 = 99,32^\circ \pm 0,04^\circ,$$

$$\theta = 17,79^\circ \pm 0,01^\circ.$$

Определим на основании этих данных удельную теплоёмкость алюминия c и её предельные ошибки, абсолютную $(\epsilon_c)_{\text{пр}}$ и относительную $(\delta_c)_{\text{пр}}$.

Обозначая разности температур $(\theta - t_0)$ и $(t_1 - \theta)$ соответственно через τ_0 и τ_1 , находим их значения, а также значения абсолютных и относительных ошибок величин M , m , τ_0 и τ_1 , причём при определении ϵ_{τ_0} и ϵ_{τ_1} применяем формулу (11), (стр. 63):

$$\tau_0 = (\theta - t_0) = 4,27^\circ,$$

$$\tau_1 = (t_1 - \theta) = 81,53^\circ,$$

$$\epsilon_M = \pm 0,2 \text{ г}, \quad \delta_M = 0,0008 = \pm 0,08\%.$$

$$\epsilon_m = \pm 0,02 \text{ г}, \quad \delta_m = 0,0003 = \pm 0,03\%,$$

$$\epsilon_{\tau_0} = \pm 0,02^\circ, \quad \delta_{\tau_0} = 0,0047 = \pm 0,47\%,$$

$$\epsilon_{\tau_1} = \pm 0,05^\circ, \quad \delta_{\tau_1} = 0,0006 = \pm 0,06\%.$$

Для определения предельной абсолютной ошибки c находим частные дифференциалы функции (12) по всем независимым переменным, т. е. по M , m , τ_0 и τ_1 :

$$\partial c_M = \frac{\tau_0}{m\tau_1} dM,$$

$$\partial c_{\tau_0} = \frac{M}{m\tau_1} d\tau_0,$$

$$\partial c_m = -\frac{M\tau_0}{m^2\tau_1} dm,$$

$$\partial c_{\tau_1} = -\frac{M\tau_0}{m\tau_1^2} d\tau_1.$$

Применяя формулу (5'), находим $(\epsilon_c)_{\text{пр}}$:

$$(\epsilon_c)_{\text{пр}} = \pm \left(\frac{\tau_0}{m\tau_1} \epsilon_M + \frac{M}{m\tau_1} \epsilon_{\tau_0} + \frac{M\tau_0}{m^2\tau_1} \epsilon_m + \frac{M\tau_0}{m\tau_1^2} \epsilon_{\tau_1} \right).$$

Подставляя числовые данные, имеем:

$$(\epsilon_c)_{\text{пр}} = \pm \left(\frac{4,27^\circ}{62,31 \text{ г} \cdot 81,53^\circ} \cdot 0,2 \text{ г} + \frac{250 \text{ г}}{62,31 \text{ г} \cdot 81,53^\circ} \cdot 0,02^\circ + \right. \\ \left. + \frac{250 \text{ г} \cdot 4,27^\circ}{(62,31)^2 \text{ г}^2 \cdot 81,53^\circ} \cdot 0,02 \text{ г} + \right. \\ \left. + \frac{250 \text{ г} \cdot 4,27^\circ}{62,31 \text{ г} (81,53^\circ)^2} \cdot 0,05^\circ \right) \text{ кал г}^{-1} \text{ град}^{-1}.$$

Производя вычисления, получаем:

$$(\epsilon_c)_{\text{пр}} = \pm (0,00017 + 0,001 + 0,00007 + \\ + 0,00013) \text{ кал г}^{-1} \text{ град}^{-1} = \pm 0,001 \text{ кал г}^{-1} \text{ град}^{-1}.$$

Приближённое значение теплоёмкости алюминия c находим из формулы (12):

$$c = \frac{250 \text{ г} \cdot 4,27^\circ}{62,31 \text{ г} \cdot 81,53^\circ} \text{ кал г}^{-1} \text{ град}^{-1} = 0,2101 \text{ кал г}^{-1} \text{ град}^{-1}.$$

Этот результат в соответствии со значением $(\epsilon_c)_{\text{пр}}$ следует написать в таком виде:

$$c = (0,210 \pm 0,001) \text{ кал г}^{-1} \text{ град}^{-1}.$$

Для определения предельной относительной ошибки применяем формулу (6'):

$$(\delta_c)_{\text{пр}} = \pm d(\ln M + \ln \tau_0 + \ln m + \ln \tau_1)$$

или

$$(\delta_c)_{\text{пр}} = \pm \left(\frac{\epsilon_M}{M} + \frac{\epsilon_{\tau_0}}{\tau_0} + \frac{\epsilon_m}{m} + \frac{\epsilon_{\tau_1}}{\tau_1} \right) = \pm (\delta_M + \delta_{\tau_0} + \delta_m + \delta_{\tau_1}).$$

Подставляя числовые данные, имеем:

$$(\delta_c)_{\text{пр}} = \pm (0,08 + 0,03 + 0,47 + 0,06) \% = \pm 0,6 \%.$$

Этот результат показывает, что предельная относительная ошибка c является большой и что наибольшее влияние на величину $(\delta_c)_{\text{пр}}$ оказывает ошибка τ_0 , т. е. разности $(\theta - t_0)$. Это объясняется тем, что разность эта

является малой величиной, вследствие чего её предельная относительная ошибка заметно возрастает (стр. 64).

Пример 2. Определение коэффициентов вязкости η жидкостей методом Стокса приводит к формуле

$$\eta = \frac{2}{9} g \frac{r^2 (d_0 - d)}{l} \cdot t, \quad (13)$$

где g обозначает ускорение силы тяжести, r , d_0 и d обозначают соответственно радиус шарика, его плотность и плотность жидкости, а l обозначает расстояние, которое проходит шарик в жидкости за время, равное t .

Обыкновенно при таких измерениях применяются очень маленькие капельки ртути, так что величину d_0 в этой формуле можно считать известной; значения g и плотности жидкости d также обыкновенно берут из таблиц. Таким образом, приходится измерять только величины r , l и t , которые и являются в данном случае независимыми переменными. При одном из таких определений η для водного раствора глицерина были получены следующие значения величин r , l и t :

$$r = (0,0112 \pm 0,0001) \text{ см},$$

$$l = (31,23 \pm 0,05) \text{ см},$$

$$t = (62,1 \pm 0,2) \text{ сек}.$$

Определим значение η и найдём его предельные ошибки, абсолютную $(\varepsilon_\eta)_{\text{пр}}$ и относительную $(\delta_\eta)_{\text{пр}}$.

Величины g , d_0 и d берём из таблиц, принимая во внимание географическую широту места наблюдения, температуру раствора глицерина и его концентрацию:

$$g = 980,1 \text{ см сек}^{-2},$$

$$d_0 = 13,55 \text{ см}^{-3} \text{ г},$$

$$d = 1,28 \text{ см}^{-3} \text{ г}.$$

Значения этих величин считаем точными числами (стр. 31).

Ошибки величин r , l и t таковы:

$$\varepsilon_r = \pm 0,0001 \text{ см}, \quad \delta_r = \pm 0,009 = \pm 0,9\%,$$

$$\varepsilon_l = \pm 0,05 \text{ см}, \quad \delta_l = \pm 0,0016 = \pm 0,16\%,$$

$$\varepsilon_t = \pm 0,2 \text{ сек}, \quad \delta_t = \pm 0,003 = \pm 0,3\%.$$

Находим предельную абсолютную ошибку η .

Для этого вычисляем частные дифференциалы функции (13) по r , l и t и берём сумму их абсолютных значений. Получаем:

$$(\varepsilon_\eta)_{\text{пр}} = \pm \frac{2}{9} g (d_0 - d) \cdot \left[\frac{2rt}{l} \varepsilon_r + \frac{r^2 t}{l^2} \varepsilon_l + \frac{r^2}{l} \varepsilon_t \right].$$

Подставляя числовые данные, имеем:

$$(\varepsilon_\eta)_{\text{пр}} = \pm \frac{2}{9} \cdot 980,1 \cdot 12,27 \frac{\text{г}}{\text{см}^3 \text{сек}^2} \left[\frac{2 \cdot 0,0112 \cdot 62,1}{31,23} \cdot 0,0001 + \right. \\ \left. + \frac{(0,0112)^2 \cdot 62,1}{(31,23)^2} \cdot 0,05 + \frac{(0,0112)^2}{31,23} \cdot 0,2 \right] \text{ см сек}.$$

Производя вычисления, получаем:

$$(\varepsilon_\eta)_{\text{пр}} = \pm 0,015 \text{ см}^{-1} \text{ г сек}^{-1}.$$

Находим приближённое значение η по формуле (13):

$$\eta = \frac{2}{9} \cdot 980,1 \cdot 12,27 \cdot \frac{(0,0112)^2}{31,23} \cdot 62,1 \text{ см}^{-1} \text{ г сек}^{-1}.$$

Вычисляя, получаем:

$$\eta = 0,6644 \text{ см}^{-1} \text{ г сек}^{-1}.$$

В этом числе в соответствии со значением $(\varepsilon_\eta)_{\text{пр}}$ следует отбросить два последних знака, так как они являются неверными. Вследствие этого окончательный результат измерений пишем в таком виде:

$$\eta = (0,66 \pm 0,01 (5)) \text{ см}^{-1} \text{ г сек}^{-1}.$$

Предельную относительную ошибку η находим, применяя формулу (6'). Так как величины g , d_0 и d мы считаем точными, то после дифференцирования получаем:

$$(\delta_\eta)_{\text{пр}} = \pm \left(2 \cdot \frac{dr}{r} + \frac{dl}{l} + \frac{dt}{t} \right) = \pm (2\delta_r + \delta_l + \delta_t).$$

Подставляя числовые величины, находим:

$$(\delta_\eta)_{\text{пр}} = \pm (1,8 + 0,16 + 0,3)\% = \pm 2,3\%.$$

Отсюда мы видим, что точность определения коэффициента вязкости η в данном случае является малоудовлетворительной: относительная ошибка η превышает 2%, причём наибольшее влияние на ошибку результата оказывает ошибка, допущенная при измерении радиуса шарика r . Это объясняется как тем, что относительная ошибка r вообще оказалась большой, так и тем, что r входит в формуле (13) во второй степени, и его ошибка в окончательном результате удваивается.

Во всех примерах, которые были рассмотрены в этом и в предыдущем параграфах, мы вычисляли абсолютные и относительные ошибки функций в полном соответствии с формулами, к которым приводит теория ошибок. При этом мы иногда встречались с необходимостью вычислять достаточно сложные строки. Однако последнее далеко не всегда является обязательным, и на практике все вычисления ошибок функций можно делать значительно проще, пользуясь следующим приёмом.

Как было отмечено выше, вычисление относительных ошибок функций в большинстве случаев является сравнительно несложной операцией. Это подтверждается и на двух только что рассмотренных примерах. Исключение здесь представляют только тригонометрические функции, где необходимо отдельно для каждой функции находить величину её относительной ошибки. В случаях же обычных алгебраических дробей, подобных выражениям (12) и (13), для вычисления их относительных ошибок достаточно найти сумму абсолютных величин относительных ошибок всех аргументов, причём, если аргумент входит в исходную формулу в степени, отличной от единицы, то его относительную ошибку необходимо предварительно умножить на показатель степени. Таким образом, мы можем упростить вычисление ошибок функций, если по формулам теории ошибок будем находить только относительные ошибки функций, применяя соответственно формулы (4), (6) и (6'). Что же касается абсолютных ошибок функций, то их можно затем определять из общего выражения абсолютной ошибки как произведения измеряемой величины на её относительную ошибку (стр. 27). Этот упрощённый приём, которым обычно и пользуются

на практике, мы будем применять в дальнейших примерах.

Пример 3. Вторым методом измерения коэффициентов вязкости η жидкостей при помощи капиллярных вискозиметров приводит к формуле

$$\eta = \frac{\pi \cdot p \cdot r^4 \cdot t}{8Vl}, \quad (14)$$

где p обозначает давление, под которым происходит истечение жидкости из капилляра, r и l обозначают радиус и длину капилляра, а V обозначает объём жидкости, который вытекает в течение времени, равного t .

При одном из определений этим методом коэффициента вязкости водного раствора глицерина небольшой концентрации были получены для величин p , r , l , V и t следующие значения:

$$\begin{aligned} p &= (20,12 \pm 0,01) \text{ мм рт. ст.} & (\delta_p &= \pm 0,0005), \\ r &= (0,570 \pm 0,003) \text{ мм} & (\delta_r &= \pm 0,005), \\ l &= (10,526 \pm 0,005) \text{ см} & (\delta_l &= \pm 0,0005), \\ V &= (5,025 \pm 0,001) \text{ см}^3 & (\delta_V &= \pm 0,0002), \\ t &= (27,34 \pm 0,02) \text{ сек} & (\delta_t &= \pm 0,0007). \end{aligned}$$

Найдём значение коэффициента вязкости η и его предельные ошибки, абсолютную $(\epsilon_\eta)_{\text{пр}}$ и относительную $(\delta_\eta)_{\text{пр}}$.

Находим предельную относительную ошибку η .

Согласно формуле (6') и тому, что было только что сказано, имеем:

$$\begin{aligned} (\delta_\eta)_{\text{пр}} &= \pm \left(\frac{dp}{p} + 4 \frac{dr}{r} + \frac{dl}{l} + \frac{dV}{V} + \frac{dt}{t} \right) = \\ &= \pm (\delta_p + 4\delta_r + \delta_l + \delta_V + \delta_t). \end{aligned}$$

Подставляя числовые данные, находим:

$$\begin{aligned} (\delta_\eta)_{\text{пр}} &= \pm (0,0005 + 0,020 + 0,0005 + 0,0002 + 0,0007) = \\ &= \pm 0,0219 = \pm 2,2\%. \end{aligned}$$

Точность определения коэффициента вязкости в данном случае оказывается очень невысокой: относительная ошибка η превышает 2%, причём эта ошибка почти

полностью зависит от ошибки, допущенной при измерении r . Радиус капилляра трудно поддаётся точному измерению и одновременно в формуле (14) он входит в четвёртой степени; вследствие этого относительная ошибка величины r^4 оказывается очень большой.

Значение η вычисляем по формуле (14), причём p необходимо выразить в абсолютных единицах давления. Имеем:

$$\eta = \frac{20,12}{760} \cdot 1,01325 \cdot 10^6 \frac{г}{см сек^2} \cdot \frac{\pi \cdot (0,057)^4 см^4 \cdot 27,34 сек}{8 \cdot 5,025 см^3 \cdot 10,526 см}.$$

Так как относительная ошибка η оказалась очень большой, то при вычислении этой строки можно ограничиться небольшой точностью, около 1%. В результате вычисления находим:

$$\eta = 0,0574 см^{-1} г сек^{-1}.$$

Находим предельную абсолютную ошибку η :

$$(\epsilon_\eta)_{пр} = \pm 0,022 \cdot 0,0574 см^{-1} г сек^{-1} = \pm 0,001 см^{-1} г сек^{-1}.$$

Окончательное значение η следует написать в таком виде:

$$\eta = (0,057 \pm 0,001) см^{-1} г сек^{-1},$$

так как четвёртый десятичный знак в данном случае является лишним.

Пример 4. По современным представлениям квантовой механики движение электрона может быть описано некоторым волновым процессом, причём длина соответствующей волны λ определяется формулой

$$\lambda = \frac{h}{mv}, \quad (15)$$

где h , m и v обозначают соответственно постоянную Планка, массу движущегося электрона и скорость его движения.

Последняя величина зависит от массы электрона m , от его заряда e и от той разности потенциалов V , под действием которой электрон приобрёл скорость v . Эта

зависимость определяется таким выражением:

$$v = \sqrt{\frac{2eV}{m}}.$$

Вследствие этого предыдущую формулу можно написать в таком виде:

$$\lambda = \frac{h}{\sqrt{2meV}}. \quad (16)$$

Попробуем подсчитать, с какой точностью можно определить из этой формулы эквивалентную длину волны λ , если воспользоваться значениями h , m и e , принятыми в настоящее время (стр. 37 и 44). По отношению к уско-ряющему потенциалу V сделаем предположение, что его значение невелико, порядка 100—200 в, так что поправку на изменение массы от скорости можно не вводить, и что значение V определено очень точно с относительной ошибкой, равной 0,01%.

Значение предельной относительной ошибки λ находим по формуле (6'):

$$(\delta_\lambda)_{пр} = \pm \left(\delta_h + \frac{1}{2} \delta_m + \frac{1}{2} \delta_e + \frac{1}{2} \delta_V \right).$$

Из таблиц находим, что относительные ошибки h , m и e определяются такими числами:

$$\delta_h = 0,15\%,$$

$$\delta_m = 0,11\%,$$

$$\delta_e = 0,01\%.$$

Вследствие этого имеем:

$$(\delta_\lambda)_{пр} = \pm (0,15 + 0,06 + 0,005 + 0,005)\% = 0,22\%.$$

Мы видим, что, пользуясь формулой (16), длину волны λ можно определить с точностью до 0,2%, причём наибольшее влияние на ошибку λ оказывает ошибка h , ошибки e и V роли не играют, т. е. эти величины в данном случае можно считать точными,

4. Вторая задача теории ошибок. Второй, или обратной, задачей теории ошибок является определение ошибок аргументов, если известны ошибки функции и вид функциональной зависимости. Таким образом, в основном уравнении

$$y \pm dy = f(x_1 \pm dx_1, x_2 \pm dx_2, x_3 \pm dx_3, \dots, x_n \pm dx_n)$$

ошибка dy считается известной, а ошибки $dx_1, dx_2, dx_3, \dots, dx_n$ мы должны определить. Как уже было сказано, с такими задачами мы встречаемся в тех случаях, когда необходимо найти значение некоторой величины y с определённой точностью, заранее установленной. Удовлетворить этому условию мы сумеем только в том случае, если заранее, ещё не приступая к измерениям величин $x_1, x_2, x_3, \dots, x_n$, определим их возможные ошибки, т. е. такие ошибки этих величин, при которых ошибка y не может выйти за намеченные пределы.

Эта задача в общем случае оказывается неопределённой, что нетрудно видеть из следующих соображений.

Общее выражение предельной относительной ошибки функции f , как это следует из формулы (6'), должно иметь вид многочлена с n членами, каждый из которых в общем случае представляет собой произведение относительной ошибки одного из аргументов на некоторый коэффициент. Таким образом, можно написать:

$$(\delta_y)_{\text{пр}} = \pm \left(a_1 \frac{dx_1}{x_1} + a_2 \frac{dx_2}{x_2} + a_3 \frac{dx_3}{x_3} + \dots + a_n \frac{dx_n}{x_n} \right), \quad (17)$$

где коэффициенты $a_1, a_2, a_3, \dots, a_n$ являются показателями степени соответствующих аргументов, как это имело место, например, при определении относительных ошибок функций (13), (14) и (16); кроме того, члены этого многочлена могут содержать тригонометрические величины, как это мы видели, например, при определении относительных ошибок функции (11).

Из последнего равенства непосредственно видно, что задача о нахождении ошибок аргументов $x_1, x_2, x_3, \dots, x_n$ по данным ошибкам функции y в общем случае является неопределённой, так как, очевидно, удовлетворить усло-

вию (17) можно при чрезвычайно различных значениях ошибок аргументов $dx_1, dx_2, dx_3, \dots, dx_n$.

Однако на практике для этой задачи в большинстве случаев удаётся находить удовлетворительное решение, хотя иногда и не вполне определённое. Эти вопросы можно разъяснить следующим образом.

Относительно многочлена (17) можно, прежде всего, предположить, что все его члены должны оказывать одинаковое влияние на ошибку функции y . Это предположение очень часто называют принципом или методом равных влияний. При таком условии мы можем написать:

$$a_1 \delta x_1 = a_2 \delta x_2 = a_3 \delta x_3 = \dots = a_n \delta x_n = \frac{(\delta_y)_{\text{пр}}}{n}.$$

В результате мы получаем n уравнений:

$$\delta x_1 = \pm \frac{(\delta_y)_{\text{пр}}}{na_1}; \quad \delta x_2 = \pm \frac{(\delta_y)_{\text{пр}}}{na_2}; \quad \delta x_3 = \pm \frac{(\delta_y)_{\text{пр}}}{na_3}; \quad \dots$$

$$\dots; \quad \delta x_n = \pm \frac{(\delta_y)_{\text{пр}}}{na_n}, \quad (18)$$

из которых можно определить значения относительных ошибок всех аргументов.

Этот приём очень часто приводит к положительным результатам, т. е., определив из уравнений (18) ошибки $\delta x_1, \delta x_2, \delta x_3, \dots, \delta x_n$, мы находим, что величина ошибок всех x лежит в пределах точности, доступной при их измерениях, т. е., приступая к измерениям величин $x_1, x_2, x_3, \dots, x_n$, мы можем все измерения произвести так, что их ошибки будут соответствовать условиям (18). В этих случаях поставленная задача, как мы видим, получает вполне удовлетворительное решение.

Но часто мы встречаем и такие случаи, когда полученные нами из условий (18) ошибки δx , если и не все, то некоторые из них, оказываются очень малыми, и осуществить при измерениях такую точность обычными приёмами мы не в состоянии. В результате на практике мы неизбежно встретимся с возрастанием соответствующих членов в выражении (17), а значит, и с возрастанием $(\delta_y)_{\text{пр}}$. Чтобы избежать последнего, мы могли бы

попытаться уменьшить ошибки остальных x , но это, очевидно, далеко не всегда окажется возможным. Вследствие этого мы встречаемся с необходимостью или всё же повысить точность измерения тех x , которые вызывают возрастание ошибки функции y , или же понизить наши начальные требования к точности её определения. Мы видим, что в этих случаях решение поставленной задачи носит далеко не определённый характер, и успех здесь, очевидно, во многом зависит от искусства и опытности экспериментатора.

Одновременно необходимо обратить внимание на следующее обстоятельство.

Исследование, подобное только что описанному, можно применить к любой формуле, которая служит для определения одной какой-либо величины по данным значениям других величин, и для каждой формулы можно составить выражение (17). Теоретически для целей дальнейшего исследования мы считаем, что все члены в правой части этого уравнения имеют одинаковую величину, и отсюда находим по уравнениям (18) относительные ошибки всех x . Но одновременно мы можем вычислить величину отдельных членов в уравнении (17), принимая во внимание предельную точность, доступную на практике при измерениях каждого из x . При этом очень часто мы видим, что некоторые из членов этого уравнения оказываются значительно бóльшей величины, чем остальные члены. Это говорит или о том, что среди x есть такие физические величины, при измерении которых не удаётся получить большой точности, или же о том, что некоторые из x входят в основную формулу в степени с большим показателем, например в третьей или четвёртой степени. В подобных случаях мы имеем право считать, что данный метод измерения имеет принципиальный недостаток. Примерами подобного рода измерительных методов с принципиальными недостатками могут служить два рассмотренных нами метода определения коэффициентов вязкости жидкостей: метод Стокса и метод капиллярного вискозиметра, представленные формулами (13) и (14). В первой из этих формул радиус шарика, а во второй внутренний радиус капилляра трудно поддаются точным измерениям; их

относительные ошибки всегда оказываются большими, а между тем в формулу (13) радиус шарика входит во второй степени, а в формулу (14) радиус капилляра входит в четвёртой степени. Можно заранее, ещё не приступая к измерениям, утверждать, что для точных измерений коэффициентов вязкости эти методы не следует применять, если не внести в них каких-либо существенных улучшений, что возможно достичь различными приёмами. Так, например, для того, чтобы исключить влияние ошибки при измерении радиуса капилляра в методе капиллярного вискозиметра, можно этим методом пользоваться только для относительных измерений. Для этого через один и тот же капиллярный вискозиметр при одинаковых условиях пропускают последовательно одинаковые объёмы двух различных жидкостей и находят отношение их коэффициентов вязкости η_0 и η_1 , которое в этом случае определяется выражением

$$\frac{\eta_1}{\eta_0} = \frac{t_1}{t_0} \cdot \frac{p_1}{p_0}.$$

Если коэффициент вязкости одной из жидкостей, например η_0 , точно известен, то последняя формула даёт возможность находить коэффициент вязкости η_1 другой жидкости, не прибегая к измерению радиуса капилляра.

Отсюда следует, что возможности, которыми располагает теория ошибок, явно выходят за пределы тех основных задач, которые были указаны раньше (стр. 84). Мы видим, что теория ошибок позволяет производить анализ, или критическую оценку, измерительных методов и выявлять их недостатки. В частности, если для определения некоторой величины y требуется измерить ряд величин, $x_1, x_2, x_3, \dots, x_n$, и мы решаем повысить точность определения y , то мы можем выяснить, какие из величин x для этого следует измерить с особенно большой точностью. Если подобный анализ мы произведём над всеми методами, которые применяются для определения какой-либо одной физической величины, например коэффициента вязкости, коэффициента преломления, заряда электрона и т. п., то мы можем выделить лучшие методы для измерения

данной величины и определить, в каком направлении следует вести работы по их дальнейшему улучшению.

Для того чтобы сделать более ясными различные вопросы, которые здесь могут возникнуть, рассмотрим несколько примеров.

Пример 1. При определении объёма цилиндра V мы измеряем высоту цилиндра h и радиус его основания r , а затем применяем обычную формулу:

$$V = \pi r^2 h.$$

Объём V следует определить с небольшой точностью так, чтобы его относительная ошибка δ_V была равна 1%. Найдём, с какой точностью следует для этого измерить величины r и h . Их приближённые значения примем равными соответственно 1 см и 5 см.

На основании формулы (6), принимая во внимание значение δ_V , имеем:

$$\begin{aligned} \delta_V &= \pm d \ln \pi r^2 h = \pm \left(2 \frac{dr}{r} + \frac{dh}{h} \right) = \\ &= \pm (2\delta_r + \delta_h) = 1\% = 0,01. \end{aligned}$$

Отсюда на основании метода равных влияний находим:

$$2\delta_r = \delta_h = \pm \frac{1}{2} \cdot 0,01 = \pm 0,005,$$

или

$$\delta_r = \pm 0,0025 = 0,25\%,$$

$$\delta_h = \pm 0,005 = 0,5\%.$$

Находим абсолютные ошибки r и h :

$$\epsilon_r = \pm 0,0025 \cdot 1 \text{ см} = \pm 0,0025 \text{ мм},$$

$$\epsilon_h = \pm 0,005 \cdot 5 \text{ см} = \pm 0,025 \text{ мм}.$$

Мы видим, что при данных размерах цилиндра радиус его основания необходимо измерить с большей точностью; абсолютная ошибка радиуса не должна превышать 0,025 мм, тогда как абсолютная ошибка h может быть в десять раз больше. Отсюда следует, что при измерении h можно воспользоваться обычным штангенциркулем,

тогда как для измерения r надёжнее применить винтовой микрометр.

По отношению к этому примеру необходимо сделать следующее весьма важное замечание.

При тех вычислениях, которые мы только что проделали, мы не приняли во внимание ошибку числа π и её влияние на ошибку V , иначе говоря, мы рассматривали π как точное число (стр. 30). Это нас обязывает, определяя значение V , взять π с таким числом десятичных знаков, чтобы его относительная ошибка была малой по сравнению с относительными ошибками r и h . Для этого π придётся взять с тремя или лучше с четырьмя десятичными знаками (стр. 25), что осложняет арифметические вычисления. Мы можем этого избежать, если весь расчёт ошибок проведём несколько иначе. Для этого нашу основную задачу мы поставим в таком виде: найдём, с какой точностью следует измерить величины r и h и с каким числом десятичных знаков следует взять π для того, чтобы относительная ошибка V не превышала 1%.

Таким образом, мы будем считать π , подобно r и h , независимой переменной, т. е. будем V считать функцией трёх величин: π , r и h . Вследствие этого, применяя формулу (6') и принимая во внимание значение δ_V , имеем:

$$\delta_V = \pm (\delta_\pi + 2\delta_r + \delta_h) = \pm 1\% = \pm 0,01.$$

Вновь применяя метод равных влияний в этом случае к трём слагаемым, находим:

$$\delta_\pi = 2\delta_r = \delta_h = \pm 0,003$$

или

$$\delta_\pi = \pm 0,003,$$

$$\delta_r = \pm 0,0015,$$

$$\delta_h = \pm 0,003.$$

Находим значения абсолютных ошибок r и h для этого случая:

$$\epsilon_r = \pm 0,0015 \cdot 1 \text{ см} = \pm 0,0015 \text{ мм},$$

$$\epsilon_h = \pm 0,003 \cdot 5 \text{ см} = \pm 0,015 \text{ мм}.$$

Из этих результатов мы видим, что π можно взять с двумя десятичными знаками ($\delta_\pi = 0,05\%$, стр. 25), что упрощает арифметические вычисления, но при этом необходимо повысить точность измерений r и h . Последнее в данном примере не вызывает каких-либо осложнений, так как новая величина абсолютных ошибок ϵ_r и ϵ_h показывает, что r и h можно измерить так же, как и прежде: h при помощи штангенциркуля и r при помощи винтового микрометра.

К такому приёму иногда прибегают, если в основную формулу входят коэффициенты, которые можно брать с различной точностью. Но при этом далеко не всегда мы получаем такой же простой результат, как в рассмотренном примере, и иногда оказывается, что повышение точности измерений, которое неизбежно при этом имеет место, может создавать большие осложнения. В этих случаях постоянные коэффициенты следует брать с той точностью, при которой они не изменяют заметно точности окончательного результата вычислений.

Пример 2. Необходимо определить удельный вес γ тела, вырезанного в форме прямоугольного параллелепипеда. Удельный вес следует определить с большой точностью, так, чтобы его относительная ошибка δ_γ не превышала $0,02\%$. Для этого мы можем определить вес тела в пустоте G и его объём V , который мы можем найти, измерив длину параллелепипеда l , его ширину d и высоту h .

Найдём величину относительных ошибок, которые мы можем допустить при взвешивании тела и при измерении l , d и h , не понижая точности определения удельного веса. Для вычисления γ пользуемся обычной формулой:

$$\gamma = \frac{G}{V} = \frac{G}{ldh}.$$

На основании формулы (6'), принимая во внимание значение δ_γ , имеем:

$$\delta_\gamma = \pm (\delta_G + \delta_l + \delta_d + \delta_h) = \pm 0,02\%.$$

Отсюда, применяя метод равных влияний, находим:

$$\delta_G = \delta_l = \delta_d = \delta_h = \pm \frac{1}{4} \cdot 0,02\% = \pm 0,005\% = \pm 0,00005.$$

Определить вес тела G с относительной ошибкой, равной $0,005\%$, не представляется сложной задачей, так как взвешивание тел можно производить со значительно большей точностью. Но измерение l , d и h с относительной ошибкой, равной также $\pm 0,005\%$, представляет собой более сложную операцию и иногда может оказаться невозможным. Действительно, если l , d и h имеют относительно большую величину, например приближённо равны 3 см, $2,5$ см и 2 см, то при такой точности измерения их абсолютные ошибки будут равны соответственно:

$$\epsilon_l = \pm 0,00005 \cdot 3 \text{ см} = \pm 0,001 (5) \text{ мм},$$

$$\epsilon_d = \pm 0,00005 \cdot 2,5 \text{ см} = \pm 0,001 \text{ мм},$$

$$\epsilon_h = \pm 0,00005 \cdot 2 \text{ см} = \pm 0,001 \text{ мм}.$$

Такая точность измерений l , d и h возможна, но потребуются применение очень хорошего сферометра или зеркального толстомера.

Если высота h параллелепипеда оказывается малой по сравнению с l и d , т. е. если от параллелепипеда мы переходим к тонким четырёхугольным пластинкам, то такой метод измерения их объёма, а значит, и удельного веса оказывается явно неточным. Если, например, толщина пластинки h приближённо равна $0,5$ мм и мы должны определить её удельный вес с прежней точностью, то h необходимо измерить с абсолютной ошибкой ϵ'_h , равной

$$\epsilon'_h = \pm 0,00005 \cdot 0,5 \text{ мм} = \pm 0,000025 \text{ мм}.$$

Такую точность мы не могли бы получить в данном случае даже от лучших интерференционных методов линейных измерений, кроме того, при измерении величин l и d было бы трудно для таких пластинок сохранить прежнюю точность. В результате мы приходим к выводу, что подобный метод измерения удельного веса в случае тонких пластинок становится грубо неточным: при толщине пластинки, равной нескольким десятым миллиметра, относительная ошибка определения удельного веса пластинки может достигать одного-двух процентов. Если такая точность недостаточна, то нам придётся отказаться от подобного метода определения γ и обратиться к более

точным методам. Например, мы можем воспользоваться методом гидростатического взвешивания, при котором все измерения ограничиваются взвешиванием тела сначала в воздухе, затем в дистиллированной воде. Нетрудно видеть также, что если мы имеем параллелепипед очень небольшого размера, например если длина его рёбер определяется несколькими миллиметрами, то определение его объёма из измерений l , d и h не может дать точных результатов; вследствие этого и в данном случае для определения γ следует воспользоваться методом гидростатического взвешивания.

Пример 3. Мы ставим перед собой задачу определить площадь треугольника S с большой точностью так, чтобы относительная ошибка S не превышала 0,01%. Для этого мы предполагаем измерить две стороны треугольника a и b и угол между ними C и воспользоваться формулой

$$S = \frac{1}{2} ab \sin C.$$

Найдём, какие ошибки мы можем допустить при измерениях a , b и C , чтобы относительная ошибка S не вышла за намеченные пределы. Предполагается, что треугольник тщательно вычерчен на зеркальной пластинке тонкими штрихами; приближённые значения a , b и C можно принять равными соответственно 5 см, 6 см и 45° .

Из формулы (6'), принимая во внимание значение δ_S , имеем:

$$\begin{aligned} \delta_S &= \pm d \ln \frac{1}{2} ab \cdot \sin C = \pm (\delta_a + \delta_b + \delta_{\sin C}) = \\ &= \pm 0,01\% = \pm 0,0001. \end{aligned}$$

На основании метода равных влияний находим:

$$\delta_a = \delta_b = \delta_{\sin C} = \pm \frac{1}{3} 0,0001 = \pm 0,00003.$$

Отсюда, принимая во внимание выражение относительной ошибки $\sin C$ (стр. 374), имеем:

$$\begin{aligned} \delta_a &= \pm 0,00003 = \pm 0,003\%, \\ \delta_b &= \pm 0,00003 = \pm 0,003\%, \\ \delta_{\sin C} &= \pm \epsilon_C \operatorname{ctg} C = \pm 0,00003 = \pm 0,003\%. \end{aligned}$$

Принимая во внимание значения a , b и C , находим их абсолютные ошибки, причём абсолютную ошибку C находим непосредственно из последнего равенства:

$$\begin{aligned} \epsilon_a &= \pm 0,00003 \cdot 5 \text{ см} = \pm 0,0015 \text{ мм}, \\ \epsilon_b &= \pm 0,00003 \cdot 6 \text{ см} = \pm 0,0018 \text{ мм}, \\ \epsilon_C &= \pm \frac{0,00003}{\operatorname{ctg} 45^\circ} = \pm 0,00003 \text{ или приближённо } \pm 6''. \end{aligned}$$

Измерение величин a , b и, в особенности, C с такой точностью в данном случае невозможно. Отсюда следует, что таким приёмом мы не можем определить площади треугольника с намеченной точностью. В лучшем случае этот приём определения S мог бы дать точность, в десять раз меньшую, т. е. относительная ошибка S была бы равна 0,1%. Необходимо обратить внимание также на то, что абсолютная ошибка ϵ_C уменьшается при уменьшении C . В данном случае это имеет отрицательное значение. Действительно, если бы значение C уменьшилось, например, до 10° , то при прежнем значении $\delta_{\sin C}$ абсолютная ошибки C оказалась бы равной:

$$\epsilon'_C = \pm \frac{0,00003}{\operatorname{ctg} 10^\circ} = \pm 0,0000053, \text{ или в угловых единицах } 1''.$$

В этом случае мы не могли бы определить значения S даже с пониженной точностью 0,1%, т. е. относительная ошибка S была бы больше этой величины.

Пример 4. При определении коэффициента α поверхностного натяжения жидкостей методом капиллярных трубок измеряют разность высот h жидких столбиков в двух вертикально установленных капиллярах, внутренние радиусы которых r_1 и r_2 предварительно тщательно измеряют. Из этих данных коэффициент поверхностного натяжения α исследуемой жидкости определяется по формуле

$$\alpha = \frac{r_1 \cdot r_2}{2(r_1 - r_2)} \gamma h,$$

где γ обозначает удельный вес жидкости.

Исследуем точность этого метода, например выясним, возможно ли этим методом определить коэффициент

поверхностного натяжения воды с относительной ошибкой δ_α не более 0,1%. Величина γ известна для воды с очень большой точностью; вследствие этого влияние ошибки γ на результаты наших вычислений можно не принимать во внимание (стр. 31). Приближённые значения величин r_1 , r_2 и h примем равными:

$$r_1 = 0,5 \text{ мм}; \quad r_2 = 0,2 \text{ мм}; \quad h = 4,5 \text{ мм}.$$

Находим ошибки, которые можно допустить при измерении r_1 , r_2 и h при условии, что

$$\delta_\alpha = 0,1\%.$$

Из формулы (6'), принимая во внимание значение δ_α , имеем:

$$\delta_\alpha = \pm (\delta_{r_1} + \delta_{r_2} + \delta_h + \delta_{(r_1-r_2)}) = \pm 0,1\% = \pm 0,001.$$

На основании принципа равных влияний находим:

$$\delta_{r_1} = \delta_{r_2} = \delta_h = \delta_{(r_1-r_2)} = \pm 0,00025.$$

Принимая во внимание значения r_1 , r_2 и h , находим их абсолютные ошибки:

$$\epsilon_{r_1} = \pm 0,00025 \cdot 0,5 \text{ мм} = \pm 0,000125 \text{ мм},$$

$$\epsilon_{r_2} = \pm 0,00025 \cdot 0,2 \text{ мм} = \pm 0,00005 \text{ мм},$$

$$\epsilon_h = \pm 0,00025 \cdot 4,5 \text{ мм} = \pm 0,01 \text{ мм}.$$

Из этих результатов мы видим, что величину h , может быть, мы сумели бы определить с такой точностью, если бы применили лучшие инструменты для линейных измерений по вертикальному направлению, но радиусы капилляров r_1 и r_2 измерить с такой точностью не представляется возможным. Даже весовой метод определения r , которым здесь обычно пользуются, мог бы дать значения r_1 и r_2 с абсолютной ошибкой, в лучшем случае равной $\pm 0,001 \text{ мм}$, т. е. относительные ошибки r_1 и r_2 оказались бы равными:

$$\delta_{r_1} = \pm 0,2\% \quad \text{и} \quad \delta_{r_2} = \pm 0,5\%.$$

Отсюда следует, во-первых, что относительная ошибка α значительно увеличится и окажется равной почти 1%

и, во-вторых, что нет оснований измерять h с той точностью, которую мы получили, и можно измерить h менее точно, например с абсолютной ошибкой, равной $\pm 0,05 \text{ мм}$, так как это вызовет лишь незначительное увеличение ошибки α .

Таким образом, наше исследование метода капиллярных трубок приводит к заключению, что этот метод не отличается большой точностью, относительная ошибка α достигает 1%, и для более точных определений коэффициента поверхностного натяжения этот метод следует считать непригодным, если не внести в него существенных изменений.

Из этих примеров становятся ясными те весьма разнообразные вопросы, которые позволяет исследовать и разрешать теория ошибок. Мы видим, что без больших трудностей можно определять как точность измерительных методов вообще, так и величину тех ошибок, которые возможны при данных условиях измерения. Мы можем также выделить методы, недостаточно точные, выяснить причины их малой точности и определить возможности их улучшения. Причины недостаточной точности измерительных методов обычно вызываются, как мы видим, или наличием среди независимых переменных величин, которые трудно поддаются точным измерениям, или неблагоприятным видом функциональной зависимости, когда, например какое-либо независимое переменное входит в функцию во второй или более высокой степени. Наличие в формуле величин, трудно поддающихся измерениям, говорит о необходимости улучшить методы измерения этих величин, а неудачный вид функциональной зависимости ставит вопрос о математическом преобразовании таких функций. В этом последнем отношении часто удаётся достигать значительных успехов. Так, например, метод капиллярного вискозиметра, теория которого приводит к формуле (14)

$$\eta = \frac{\pi \cdot p \cdot r^4 \cdot t}{8Vl},$$

является недостаточно точным вследствие наличия в этой формуле r в четвёртой степени. Оказывается, что здесь возможно математическое преобразование, которое заметно

улучшает эту формулу. Действительно, если числитель и знаменатель этой формулы умножить на длину капилляра l и принять во внимание, что $\pi r^2 l$ равняется внутреннему объёму капилляра v

$$\pi r^2 l = v,$$

то формулу (14) можно привести к виду

$$\eta = \frac{p \cdot r^2 \cdot t \cdot v}{8Vl^2}. \quad (14')$$

Таким образом, мы понизили степень r в два раза, но вместе с тем возросла на единицу степень l и появилось новое независимое переменное v . Это, однако, в данном случае оказывается менее тяжёлым, так как относительная ошибка l всегда значительно меньше относительной ошибки r (стр. 105), а v можно определить при помощи весового метода с большой точностью, так что его относительная ошибка не превысит, например, $\pm 0,03\%$. В результате этого предельная относительная ошибка η , которая раньше была равна $\pm 2,2\%$ (стр. 105), уменьшается до $\pm 1,4\%$, если при прежних исходных данных её вычислить по формуле (14'). Хотя метод продолжает оставаться недостаточно точным, однако некоторое улучшение в него мы всё же сумели внести.

5. Определение наиболее выгодных условий измерения. Нахождение лучших или наиболее выгодных условий для измерений составляет третью основную задачу теории ошибок. Такие условия, очевидно, будут иметь место в том случае, когда ошибка результата измерений окажется наименьшей. При решении этих вопросов мы вновь будем оперировать относительными ошибками функций, т. е. выражением (4) в случае одного независимого переменного и выражениями (6) и (6') в более сложных случаях двух или нескольких независимых переменных. Эти выражения являются некоторыми функциями величин x , и если мы ставим вопрос о нахождении условий, при которых относительная ошибка δ_y результата измерений является наименьшей, то с математической стороны этот вопрос решается нахождением условий минимума функций,

определяемых выражением (4) в случае одной измеряемой величины и выражениями (6) и (6') в случае двух или нескольких измеряемых величин. Отсюда мы можем сделать один очень важный вывод.

Так как минимум имеют не все функции, то не во всех случаях вопрос о нахождении наиболее выгодных условий для измерений может получить определённое решение, иначе говоря, если мы, исследуя какую-либо из функций (4) или (6), находим, что она не имеет минимума, то на вопрос о лучших условиях измерения мы принципиально не можем получить определённого ответа.

Что же касается математической стороны этого вопроса, то при нахождении минимума выражений (4) и (6) мы, очевидно, должны руководствоваться общими правилами, принятыми в дифференциальном исчислении для нахождения минимумов (и максимумов) функций одного и нескольких независимых переменных, т. е.:

а) В первом случае следует найти первую и вторую производные выражения (4) по x и, приравняв первую производную нулю, определить отсюда соответствующее значение x ; пусть оно оказалось равным a . Если при этом значении x вторая производная оказывается положительной, то функция (4) при x , равном a , имеет минимум.

б) Во втором случае следует найти частные производные первого порядка выражений (6) или (6') по всем независимым переменным и, приравняв полученные выражения нулю, определить из них значения всех переменных. Дополнительными условиями минимума вновь будет определение знака частных производных второго порядка при найденных значениях независимых переменных.

Определение лучших условий для измерений в случае нескольких независимых переменных может представлять собой сложную задачу, которая к тому же не всегда имеет определённое решение. Вследствие этого, хотя практика и выработала более простые приёмы нахождения минимумов функции нескольких независимых переменных, для решения данного вопроса к такому исследованию обращаются очень редко.

Поэтому мы подробнее исследуем только значительно более простой случай одного независимого переменного, так как здесь мы часто приходим ко вполне определённом решению в результате очень несложных математических операций. Кроме того, следует заметить, что иногда удаётся находить определённое решение значительно более простыми приёмами, даже не следуя в точности указанному выше правилу. Для того чтобы лучше разобраться в этих вопросах, рассмотрим несколько примеров.

Пример 1. В том типе гальванометров, о котором было упомянуто (стр. 90), сила тока I и угол отклонения магнитной стрелки связаны соотношением

$$I = k \operatorname{tg} \varphi.$$

Посмотрим, при каком значении угла φ чувствительность прибора будет наибольшей, т. е. мы будем производить измерение I с наименьшей относительной ошибкой.

Относительная ошибка δ_I в данном случае определяется выражением.

$$\delta_I = \pm \frac{2d\varphi}{\sin 2\varphi}.$$

Из этого выражения непосредственно видно, что δ_I получает наименьшую величину при наибольшем значении $\sin 2\varphi$, т. е. при условии, что

$$\sin 2\varphi = 1.$$

Отсюда имеем: $2\varphi = 90^\circ$ и $\varphi = 45^\circ$.

Таким образом, мы видим, что этот прибор имеет наибольшую чувствительность, если отклонения магнитной стрелки лежат в области углов, близких к 45° .

Пример 2. При фотометрических работах с поляризационным спектрофотометром установка на равенство освещения обеих половин поля прибора достигается поворотом окулярного николя. При этом сила света I_0 некоторой длины волны в спектре эталонного источника света и сила света I_x той же длины волны в спектре исследуемого источника связаны между собой формулой

$$I_x = I_0 \operatorname{tg}^2 \alpha,$$

где α обозначает угол, на который необходимо повернуть окулярный николь для установки на равенство освещения.

Определим, при каких значениях угла α чувствительность прибора оказывается наибольшей. При этом условии относительная ошибка I_x должна быть, очевидно, наименьшей.

Находим величину относительной ошибки I_x :

$$\delta_{I_x} = \pm d \ln I_0 \operatorname{tg}^2 \alpha = \pm 2d \ln \operatorname{tg} \alpha = \pm 2 \cdot \frac{2d\alpha}{\sin 2\alpha}.$$

Так же как и в предыдущем примере, минимум δ_{I_x} отвечает максимуму $\sin 2\alpha$, отсюда находим:

$$\alpha = 45^\circ,$$

т. е. наименьшую относительную ошибку результата мы будем иметь в тех случаях, когда окулярный николь приходится поворачивать на углы, близкие к 45° , иначе говоря, если при измерениях нам приходится отсчитывать углы α , сильно отличающиеся от 45° , то необходимо принимать во внимание понижение чувствительности прибора в этих областях.

Пример 3. При определении сопротивлений методом мостика измеряемое сопротивление r_x находят по формуле

$$r_x = r_0 \cdot \frac{l_1}{l_2}, \quad (19)$$

где r_0 обозначает эталонное сопротивление, величина которого очень точно измерена, а l_1 и l_2 обозначают части измерительной проволоки мостика, на которые её делит подвижной контакт при его установке на нулевое отклонение гальванометра.

Найдём наиболее выгодное положение подвижного контакта, т. е. относительную длину l_1 и l_2 , при котором ошибка r_x оказывается наименьшей.

Длину измерительной проволоки будем считать постоянной, обозначим её через L . Очевидно, что

$$L = l_1 + l_2.$$

Вследствие этого формулу (19) можно написать в таком виде:

$$r_x = r_0 \cdot \frac{l_1}{L-l_1}.$$

Находим относительную ошибку r_x :

$$\delta_{r_x} = \pm d \ln \left(r_0 \frac{l_1}{L-l_1} \right) = \pm \left(\frac{1}{l_1} + \frac{1}{L-l_1} \right) dl_1.$$

Это выражение будет иметь минимум при минимуме коэффициента при dl_1 . Обозначая его через z , т. е. полагая

$$z = \frac{1}{l_1} + \frac{1}{L-l_1}, \quad (20)$$

находим первую и вторую производные z по l_1 :

$$\begin{aligned} \frac{dz}{dl_1} &= \frac{1}{(L-l_1)^2} - \frac{1}{l_1^2}, \\ \frac{d^2z}{dl_1^2} &= \frac{2}{(L-l_1)^3} + \frac{2}{l_1^3}. \end{aligned}$$

Вторая производная при всех возможных значениях l_1 должна быть положительной, так как l_1 всегда меньше L . Поэтому функция (20) имеет минимум при том значении l_1 , которое получится, если мы положим первую производную равной нулю и решим это уравнение относительно l_1 . Вследствие этого имеем:

$$\frac{1}{(L-l_1)^2} - \frac{1}{l_1^2} = 0.$$

Отсюда находим:

$$l_1 = \frac{L}{2}.$$

Таким образом, оказывается, что наиболее выгодные условия измерения будут иметь место в том случае, если подвижной контакт мостика находится вблизи середины измерительной проволоки. Согласно формуле (19) для этого необходимо, чтобы эталонное сопротивление r_0 и измеряемое сопротивление r_x были приближённо равны по величине, что при измерениях и стараются, по возможности, осуществить.

Пример 4. Термоэлектродвижущая сила термопар при измерении её электродинамическим методом определяется по формуле

$$E = r_0 \cdot \frac{I_1 I_2}{I_1 - I_2}, \quad (21)$$

где r_0 обозначает сопротивление, введённое в цепь термопары, величина которого очень точно известна, а I_1 и I_2 обозначают силу тока в цепи термопары: I_1 — при выключенном сопротивлении r_0 , а I_2 , — если оно введено в цепь.

Посмотрим, при каком соотношении между I_1 и I_2 мы будем иметь при измерении E наименьшую относительную ошибку δ_E . Величину r_0 будем считать постоянным коэффициентом.

Рассматривая разность $I_1 - I_2$ как третье независимое переменное функции (21), применим при определении δ_E формулу (6'); получим:

$$\delta_E = \pm \left(\frac{dI_1}{I_1} + \frac{dI_2}{I_2} + \frac{d(I_1 + I_2)}{(I_1 - I_2)} \right).$$

Так как сила тока в обоих случаях измеряется одним и тем же прибором, то мы имеем право положить:

$$dI_1 = dI_2 = dI.$$

Вследствие этого выражение относительной ошибки получает такой вид:

$$\delta_E = \pm \left(\frac{1}{I_1} + \frac{1}{I_2} + \frac{2}{I_1 - I_2} \right) dI.$$

Приводя дроби в правой части к одному знаменателю, находим:

$$\delta_E = \pm \frac{I_1^2 + 2I_1 I_2 - I_2^2}{I_1^2 I_2 - I_1 I_2^2} dI.$$

Разделив числитель и знаменатель в правой части этого выражения на I_2^2 и вводя обозначение

$$\frac{I_1}{I_2} = p,$$

находим:

$$\delta_E = \pm \frac{p^2 + 2p - 1}{p - 1} \cdot \frac{dI}{I_1}. \quad (22)$$

Отсюда мы видим, что минимум δ_E должен иметь место при минимуме коэффициента при $\frac{dI}{I_1}$ в правой части этого выражения.

Обозначим этот коэффициент через q :

$$\frac{p^2 + 2p - 1}{p - 1} = q,$$

или

$$q = (p + 3) + \frac{2}{p - 1}.$$

Находим минимум этого выражения, для этого определяем его первую и вторую производные по p :

$$\frac{dq}{dp} = 1 - \frac{2}{(p - 1)^2},$$

$$\frac{d^2q}{dp^2} = \frac{4}{(p - 1)^3}.$$

Первую производную приравняем нулю и из полученного уравнения определяем p . Находим:

$$p = 1 + \sqrt{2} = 2,41.$$

При этом значении p вторая производная имеет положительное значение. Отсюда следует, что при p , равном 2,41, функция (22) имеет минимум.

Таким образом, мы вправе сказать, что наименьшую ошибку при измерении E этим методом мы будем иметь в том случае, если в цепь термопары введём такое сопротивление r_0 , которое уменьшит силу тока в 2,4 раза по сравнению с её начальным значением, т. е. при выключенном сопротивлении r_0 .

Пример 5. При определении длины волны λ световых лучей с помощью дифракционной решётки применяется формула, с которой мы уже встречались (стр. 92):

$$\lambda = \frac{1}{n} \cdot D \sin \varphi.$$

Попробуем определить те значения угла φ , при которых относительная ошибка λ будет наименьшей.

Значение относительной ошибки в данном случае определяется выражением

$$\delta_\lambda = \pm \operatorname{ctg} \varphi \cdot d\varphi.$$

Из этого выражения без дальнейших преобразований непосредственно видно, что задача о нахождении наиболее выгодных значений угла φ в данном случае не может иметь определённого решения. Действительно, минимум δ_λ должен был бы отвечать минимуму $\operatorname{ctg} \varphi$, но котангенс при изменении угла изменяется монотонно, получая нулевое значение при φ , равном 90° , что в данном случае не может представлять интереса. Таким образом, если при определении λ этим методом и можно обеспечить наиболее выгодные условия измерения, то следует иметь в виду некоторые другие факторы, например величину абсолютной ошибки φ , разрешающую силу решётки, относительную яркость спектров различного порядка и т. п.

Эти примеры показывают, что такой предварительный анализ различных измерительных методов, чрезвычайно ценный, в большинстве случаев не требует сколько-нибудь сложных расчётов. В частности, следует обратить внимание на четвёртый пример: в этом примере имеем функцию двух независимых переменных, которую вначале мы рассматривали даже как функцию трёх переменных, а при нахождении условий минимума её ошибки задачу оказалось возможным привести к случаю одного независимого переменного, и мы получили вполне определённый ответ в результате очень несложных вычислений.

ГЛАВА IV

ЗАКОН НОРМАЛЬНОГО РАСПРЕДЕЛЕНИЯ СЛУЧАЙНЫХ ОШИБОК

1. Случайные явления и их общая классификация.
Мы часто встречаемся с такими явлениями, причины возникновения которых остаются для нас неизвестными, неясными. Это объясняется тем, что подобные явления, получившие название случайных явлений, возникают в результате чрезвычайно сложной причинной цепи различных событий, проследить которую во всех деталях мы не в состоянии.

Такое элементарное определение случайности как события, причины возникновения которого остаются неизвестными, вполне достаточное в данном случае, несколько не противоречит общему философскому определению случайности как одной из форм проявления необходимости.

Можно привести очень много примеров случайных явлений. Если, например, в ящике находится некоторое количество шаров, различных по цвету, но вполне одинаковых по всем остальным признакам, и мы, не глядя, вынимаем один шар, то цвет вынутого шара будет, очевидно, случайным явлением. К числу таких же случайных явлений приходится относить и те случайные ошибки (стр. 32), которые мы неизбежно делаем при измерениях.

Иногда может казаться, что причины, вызывающие некоторое явление, нам известны, и тем не менее оно продолжает оставаться явно случайным. Так, в примере с шарами можно, конечно, утверждать, что цвет вынутого шара должен зависеть от относительного числа шаров различного цвета, находящихся в ящике; это условие

мы можем подобрать как угодно, т. е. число шаров различного цвета в ящике нам может быть известно в точности. Но ряд других условий, например расположение шаров в ящике, движения нашей руки при вынимании шара и т. д., остаётся для нас неизвестным, и всё явление продолжает носить случайный характер.

Несмотря на неопределённость причин, которые вызывают возникновение случайных явлений, всё же оказывается, что при некоторых условиях случайные явления подчиняются определённым закономерностям. Законы случайных явлений изучаются в теории вероятностей, которая имеет одну особенность: все выводы, основанные на этой теории, носят только вероятный характер, что является следствием основного свойства случайных явлений.

Различным случайным явлениям, которые в теории вероятностей принято называть **событиями**, приписывают определённые признаки, или свойства, и делят их на отдельные виды. Для дальнейшего достаточно пока указать два вида событий:

1) События называются **единственно возможными**, если при наступлении условий, необходимых для возникновения событий, одно из них непременно происходит. Так, например, если в ящике находятся шары чёрного, белого и красного цветов, то вынутый шар (условие, необходимое для наступления события) непременно будет или белого, или чёрного, или красного цвета; если мы стреляем в цель (условие, необходимое для наступления события), то мы или попадаем в неё, или делаем промах и т. д. Во всех подобных случаях мы, очевидно, имеем дело с событиями, **единственно возможными**.

2) События называются **несовместимыми**, если возникновение одного из них совершенно исключает возможность одновременного возникновения хотя бы одного из остальных. Примерами событий такого вида могут служить приведённые выше примеры единственно возможных событий, так как появление белого шара, очевидно, исключает возможность одновременного появления шара чёрного или красного, попадание в цель исключает возможность одновременного промаха и т. п.

Таким образом, мы вправе сказать, что в обоих приведённых примерах события удовлетворяют обоим признакам, т. е. эти события оказываются единственно возможными и несовместимыми.

В таких простых случаях соответствие событий с тем или другим признаком очень нетрудно установить с полной определённойностью. Но в более сложных случаях анализ общих признаков данного конкретного ряда событий часто оказывается делом далеко не лёгким; иногда оказывается также, что результаты подобного анализа приводят к выводам, не вполне определённым или не вполне убедительным. Вследствие этого, приписывая различным событиям наличие тех или иных признаков, нам часто приходится проявлять большую осторожность.

Все события могут происходить, если их возникновению благоприятствуют хотя бы некоторые из тех условий, при которых мы можем ожидать появления этих событий, или, как обыкновенно говорят, из всех возможных случаев. Например, если в ящике находятся только белые и чёрные шары и мы попрежнему вынимаем один шар, то может появиться белый или чёрный шар, причём мы вправе считать, что число случаев, которые благоприятствуют появлению белого и чёрного шаров, определяется числом шаров того и другого цвета в ящике. Одновременно мы можем считать, что общее число всех возможных случаев определяется общим числом шаров и того и другого цвета в ящике.

Рассмотрим вначале две предельные возможности.

Во-первых, допустим, что появлению некоторого события благоприятствуют все возможные случаи, например мы допускаем, что в ящике находятся только белые шары. В этом случае вынутый шар будет, конечно, белым. Такое событие, т. е. событие, которое должно произойти непременно, называется достоверным событием: все возможные случаи являются для него благоприятными.

Во-вторых, допустим, что появлению некоторого события не благоприятствует ни один из возможных случаев; например, мы допускаем, что в ящике нет ни одного

белого шара. В таком случае мы, очевидно, не можем вынуть из ящика белый шар. Такое событие, которое не может произойти ни в каком случае, называется невозможным событием: ни один из возможных случаев ему не благоприятствует.

Наконец, допустим, что некоторая часть всех возможных случаев благоприятствует появлению события, а остальные не благоприятствуют ему. В таком случае это событие мы не можем считать ни достоверным, ни невозможным, так как оно, очевидно, может или произойти, или не произойти. Такие события называют вероятными событиями: некоторые случаи благоприятствуют их появлению, а остальные не благоприятствуют.

2. Математическое определение вероятности. Несомненно, что вероятность какого угодно события должна быть тем больше, чем больше число случаев, которые ему благоприятствуют, по сравнению с общим числом всех возможных случаев. Вследствие этого соображения условились математическую вероятность события определять отношением числа n случаев, благоприятствующих событию, к общему числу N всех возможных случаев, причём при более строгом определении вероятности следует указать, что все N случаев должны быть единственно возможными и несовместимыми.

Таким образом, обозначая вероятность некоторого события через P , можно написать:

$$P = \frac{n}{N}. \quad (1)$$

Из формулы (1) мы можем сделать следующие очевидные выводы:

1) Если событие является достоверным, т. е. все возможные случаи ему благоприятствуют, то в формуле (1) надо положить:

$$n = N$$

или

$$P = 1,$$

т. е. вероятность достоверного события всегда равна единице.

2) Если событие является невозможным, т. е. ни один случай ему не благоприятствует, то в формуле (1) надо положить:

$$n = 0$$

или

$$P = 0,$$

т. е. вероятность невозможного события всегда равна нулю.

3) Если событие является возможным, но не достоверным, то одни случаи благоприятствуют его возникновению, а другие не благоприятствуют; вследствие этого в формуле (1) всегда должно быть:

$$n < N,$$

т. е. вероятность любого возможного события всегда выражается некоторой правильной дробью.

Рассмотрим какой-либо пример. Пусть, например, из полной колоды карт, хорошо перетасованных, мы, не глядя, вынимаем одну карту. Спрашивается, какова вероятность того, что вынутая карта окажется определённой, выбранной заранее масти, например бубновой. В полной колоде, как известно, 52 карты, по 13 карт каждой масти. Таким образом, число случаев, благоприятствующих появлению карты бубновой масти, равно 13, а общее число всех возможных случаев равно 52; отсюда следует, что вероятность P_1 появления карты бубновой масти равна:

$$P_1 = \frac{13}{52} = \frac{1}{4}.$$

Такова же, очевидно, и вероятность появления карты любой другой масти.

Далее, какова вероятность появления какой-либо определённой карты, например валета? Таких карт в колоде четыре; отсюда следует, что вероятность P_2 появления валета равна:

$$P_2 = \frac{4}{52} = \frac{1}{13}.$$

Наконец, какова вероятность появления определённой карты определённой масти, например валета бубен. Такая карта в колоде только одна; отсюда следует, что вероятность P_3 появления валета бубен равна:

$$P_3 = \frac{1}{52}.$$

Для дальнейшего полезно заметить, что произведение вероятности P_1 появления карты бубновой масти на вероятность P_2 появления валета оказывается равным вероятности P_3 появления валета бубновой масти. Действительно,

$$P_1 \cdot P_2 = \frac{1}{4} \cdot \frac{1}{13} = \frac{1}{52} = P_3.$$

Если вероятность некоторого события выражается правильной дробью, близкой к единице, например 0,99 или 0,999, то такое событие обыкновенно считают событием практически достоверным. Точно так же, если вероятность некоторого события выражается дробью, близкой к нулю, например 0,01 или 0,001, то такое событие обыкновенно считают событием практически невозможным.

Исходя из определения математической вероятности события, мы можем ввести понятие ещё о двух видах событий, а именно:

1) Два события мы называем **противоположными**, если число случаев, благоприятствующих возникновению и того и другого события, в сумме равно общему числу всех возможных случаев. Так, если в ящике находятся шары только белого и чёрного цветов, и мы попрежнему вынимаем один шар, то появление белого и появление чёрного шара являются событиями противоположными, так как число шаров того и другого цвета

вместе равно общему числу шаров в ящике. Из определения противоположности событий следует, что сумма вероятностей обоих противоположных событий всегда равна единице, т. е. равна достоверности.

2) События мы называем независимыми, если возникновение одного или нескольких из них не изменяет вероятности возникновения каждого из остальных.

3. Теоремы сложения и умножения вероятностей. В основе теории вероятностей лежат две теоремы, которые обыкновенно называют: 1) теоремой сложения вероятностей и 2) теоремой умножения вероятностей. Эти теоремы можно изложить следующим образом.

1. Теорема сложения вероятностей.

Если имеется несколько несовместимых событий, то вероятность возникновения одного из них без определения, какого именно, равняется сумме вероятностей этих событий.

Для доказательства этой теоремы допустим, что имеется m несовместимых событий:

$$A_1, A_2, A_3, \dots, A_m. \quad (2)$$

Число всех возможных случаев, при которых возникают события ряда (2), считаем равным N , причём:

n_1 случаев благоприятствуют событию A_1 ,

n_2 случаев благоприятствуют событию A_2 ,

n_3 случаев благоприятствуют событию A_3 ,

.....

n_m случаев благоприятствуют событию A_m .

Для общности рассуждений мы можем предположить, что в числе всех N возможных случаев есть некоторое число таких, которые не благоприятствуют ни одному из событий ряда (2), так что в общем случае:

$$n_1 + n_2 + n_3 + \dots + n_m < N. \quad (3)$$

Обозначая вероятности каждого из событий $A_1, A_2, A_3, \dots, A_m$ соответственно через $P_1, P_2, P_3, \dots, P_m$,

мы можем согласно формуле (1) написать:

$$P_1 = \frac{n_1}{N}, \quad P_2 = \frac{n_2}{N}, \quad P_3 = \frac{n_3}{N}, \quad \dots, \quad P_m = \frac{n_m}{N}. \quad (4)$$

Так как события $A_1, A_2, A_3, \dots, A_m$ мы считаем несовместимыми, то в числе случаев $n_1, n_2, n_3, \dots, n_m$ не может быть случаев одинаковых, которые благоприятствовали бы одновременному возникновению двух или нескольких событий A , — это исключено вследствие их несовместимости. Таким образом, сумма величин n состоит попрежнему из случаев несовместимых, причём эта сумма определяет число случаев, благоприятных возникновению одного какого-либо из событий $A_1, A_2, A_3, \dots, A_m$ без определения, какого именно. Вследствие этого, вычисляя вероятность P возникновения одного из событий ряда (2), безразлично, какого именно, мы можем на основании формулы (1) написать:

$$P = \frac{n_1 + n_2 + n_3 + \dots + n_m}{N} \quad (5)$$

ИЛИ

$$P = \frac{n_1}{N} + \frac{n_2}{N} + \frac{n_3}{N} + \dots + \frac{n_m}{N} .$$

Отсюда в соответствии с равенствами (4) получаем:

$$P = P_1 + P_2 + P_3 + \dots + P_m. \quad (6)$$

Из выражения (6) видно, что вероятность появления одного из нескольких несовместимых событий, без определения, какого именно, всегда больше вероятности каждого из этих событий в отдельности.

Если мы предположим, в частности, что в числе всех возможных случаев нет ни одного, который не благоприятствует ни одному из событий ряда (2), то вместо неравенства (3) мы можем написать:

$$n_1 + n_2 + n_3 + \dots + n_m = N. \quad (7)$$

На основании этого равенства из выражения (5) получаем:

$$P = 1, \quad (8)$$

т. е. вероятность возникновения одного из событий ряда (2), безразлично, какого именно, при наличии условия (7) становится достоверностью. Этот вывод вполне понятен, так как из ряда событий $A_1, A_2, A_3, \dots, A_m$ одно какое-либо непременно должно произойти, если все N возможных случаев благоприятствуют этим событиям; иначе можно сказать, что события ряда (2), которые мы вначале определили только как несовместимые, становятся при наличии условия (7), кроме того, событиями единственно возможными, т. е. одно какое-либо событие непременно должно возникнуть, что и подтверждается равенством (8).

2. Теорема умножения вероятностей.

Если имеется несколько независимых событий, то вероятность совместного возникновения этих событий равняется произведению их вероятностей.

Одновременное возникновение нескольких независимых событий рассматривается как некоторое новое событие, которое принято называть сложным событием. Таким образом, теорему умножения вероятностей можно формулировать ещё так: вероятность сложного события, которое состоит в совместном возникновении нескольких независимых событий, равняется произведению вероятностей этих событий.

Докажем эту теорему сначала для двух событий. Для этого допустим, что имеется два независимых события A_1 и A_2 . Обозначим число возможных случаев, при которых могут возникать события A_1 и A_2 , соответственно через N_1 и N_2 , а число случаев, которые благоприятствуют этим событиям, соответственно через n_1 и n_2 . Вероятности событий A_1 и A_2 обозначим через P_1 и P_2 ; на основании формулы (1) эти вероятности равны:

$$P_1 = \frac{n_1}{N_1} \text{ и } P_2 = \frac{n_2}{N_2}. \quad (9)$$

Совместное возникновение обоих событий A_1 и A_2 будем рассматривать, как некоторое новое сложное событие. Для того чтобы определить его вероятность, не-

обходимо определить: 1) число всех возможных случаев, при которых может возникать это событие, и 2) число случаев, которые ему благоприятствуют. Обозначим первое число случаев через N и второе — через n . При вычислении N необходимо каждый из N_1 возможных случаев, при которых может возникать событие A_1 , сопоставить с каждым из N_2 возможных случаев, при которых может возникать событие A_2 , так как каждое такое сопоставление даёт нам один из случаев, при котором могут возникать совместно оба события A_1 и A_2 . Таким образом, можно написать:

$$N = N_1 \cdot N_2. \quad (10)$$

Совершенно так же при вычислении n необходимо каждый из n_1 случаев, благоприятствующих событию A_1 , сопоставить с каждым из n_2 случаев, благоприятствующих событию A_2 , так как каждое такое сопоставление даёт нам один из случаев, который благоприятствует совместному возникновению обоих событий A_1 и A_2 . Таким образом, можно написать:

$$n = n_1 \cdot n_2. \quad (11)$$

Обозначая вероятность совместного возникновения обоих событий A_1 и A_2 через P , на основании формулы (1) и равенств (10) и (11) можно написать:

$$P = \frac{n_1 n_2}{N_1 N_2}$$

или

$$P = \frac{n_1}{N_1} \cdot \frac{n_2}{N_2}.$$

Из последнего выражения на основании равенств (9) находим:

$$P = P_1 \cdot P_2, \quad (12)$$

что и является доказательством теоремы умножения вероятностей для случая двух независимых событий.

Очевидно, что теорему умножения вероятностей, доказанную для случая двух независимых событий, мы можем при помощи её последовательного применения

распространить на любое число независимых событий. Таким образом, можно сказать следующее:

Если имеется m независимых событий

$$A_1, A_2, A_3, \dots, A_m,$$

вероятности которых равны соответственно $P_1, P_2, P_3, \dots, P_m$, то вероятность P совместного возникновения всех этих событий определяется произведением всех P :

$$P = P_1 \cdot P_2 \cdot P_3 \cdot \dots \cdot P_m. \quad (12')$$

Из выражения (12') мы видим, что вероятность одновременного возникновения нескольких независимых событий всегда меньше вероятности каждого из этих событий в отдельности.

Для того чтобы познакомиться с применением теорем сложения и умножения вероятностей, рассмотрим какой-либо пример. Пусть, например, в двух ящиках находятся шары, которые отличаются только цветом. В одном ящике 15 белых, 20 красных и 25 синих шаров, всего 60 шаров; в другом ящике 10 белых, 25 красных и 15 синих шаров, всего 50 шаров. Шары в ящиках тщательно перемешаны. Не глядя, вынимаем по одному шару из каждого ящика.

Спрашивается, какова вероятность того, что оба вынутых шара окажутся одинакового цвета, т. е. будет вынуто два белых, или два красных, или два синих шара.

Находим вероятности $P'_б$, $P'_к$ и $P'_с$, т. е. вероятности появления белого, красного и синего шаров из первого ящика:

$$P'_б = \frac{15}{60} = \frac{1}{4}; \quad P'_к = \frac{20}{60} = \frac{1}{3}; \quad P'_с = \frac{25}{60} = \frac{5}{12}.$$

Находим вероятности $P''_б$, $P''_к$ и $P''_с$, т. е. вероятности появления белого, красного и синего шаров из второго ящика:

$$P''_б = \frac{10}{50} = \frac{1}{5}; \quad P''_к = \frac{25}{50} = \frac{1}{2}; \quad P''_с = \frac{15}{50} = \frac{3}{10}.$$

По теореме умножения вероятностей находим вероятности $P_б$, $P_к$ и $P_с$ появления из ящиков двух белых, двух

красных и двух синих шаров:

$$P_б = P'_б \cdot P''_б = \frac{1}{4} \cdot \frac{1}{5} = \frac{1}{20},$$

$$P_к = P'_к \cdot P''_к = \frac{1}{3} \cdot \frac{1}{2} = \frac{1}{6},$$

$$P_с = P'_с \cdot P''_с = \frac{5}{12} \cdot \frac{3}{10} = \frac{1}{8}.$$

Таким образом, поставленная задача решается очень просто.

Если мы хотим узнать, какова вероятность $P_ц$ появления двух шаров одного цвета, безразлично какого, то следует применить теорему сложения вероятностей; таким образом, находим:

$$P_ц = P_б + P_к + P_с = \frac{1}{20} + \frac{1}{6} + \frac{1}{8} = \frac{41}{120},$$

т. е. немного больше одной трети.

Если мы хотим узнать, какова вероятность P появления двух шаров цветных, безразлично красных или синих, то должны сложить вероятности $P_к$ и $P_с$; получаем:

$$P = P_к + P_с = \frac{1}{6} + \frac{1}{8} = \frac{7}{24},$$

т. е. немного меньше одной трети.

Интересно отметить, что вероятность появления двух шаров одного цвета, т. е. $P_б$, $P_к$ или $P_с$, оказывается меньшей по сравнению с вероятностью появления шаров того же цвета из каждого ящика в отдельности, что находится в полном соответствии с замечанием, сделанным при выводе теоремы умножения вероятностей. Одновременно вероятность появления двух шаров одинакового цвета, безразлично какого, т. е. $P_ц$, оказывается больше каждой из вероятностей появления двух шаров одного цвета, что находится в полном соответствии с замечанием, сделанным при выводе теоремы сложения вероятностей.

4. Общие закономерности случайных ошибок. Основные положения теории вероятностей, рассмотренные выше, мы применяем к событиям единственно возможным, несовместимым и независимым. Пользуясь этими положениями

в теории случайных ошибок, мы должны рассматривать последние так же, как события единственно возможные, несовместимые и независимые.

Имеются все основания считать, что случайные ошибки удовлетворяют этим требованиям. Действительно,

1) Производя ряд измерений какой-либо величины, мы при каждом из них неизбежно делаем некоторую случайную ошибку. Отсюда следует, что случайные ошибки являются событиями единственно возможными.

2) При каждом отдельном измерении мы можем сделать только одну случайную ошибку, так как одновременно две случайные ошибки при одном отдельном измерении появиться, очевидно, не могут. Отсюда следует, что случайные ошибки являются событиями несовместимыми.

3) Случайные ошибки являются, очевидно, событиями независимыми, так как появление одной какой-либо случайной ошибки не изменяет вероятности появления всех остальных случайных ошибок.

Таким образом, мы видим, что при разработке теории случайных ошибок можно применять положения теории вероятностей, однако при этом безусловно необходимо иметь в виду ещё следующее:

1) В теории случайных ошибок предполагается, что все измерения данного ряда выполнены с одинаковой точностью, т. е. являются измерениями равноточными и вследствие этого равноценными, поэтому все измерения необходимо делать одинаково тщательно, так как иначе кроме случайных ошибок, которые лежат в пределах точности измерений, в них могут оказаться ещё более грубые ошибки, обусловленные нашей небрежностью.

2) Выводы теории вероятностей оправдываются только при многократном повторении событий. Поэтому вся теория ошибок разработана в предположении, что измерения всегда повторяются несколько раз, точнее говоря, много раз. Отсюда следует, что выводы теории ошибок мы имеем право применять только к тем измерениям, которые удовлетворяют этому требованию, т. е. были повторены достаточно большое число раз. Только при этом условии окончательный результат измерений, вычисленный на

основании положений теории случайных ошибок, может отвечать в пределах точности измерений значению измеряемой величины.

3) Теория случайных ошибок, как основанная полностью на теории вероятностей, даёт возможность установить в конечном результате только наиболее вероятное значение измеряемой величины. Тем не менее, при обработке результатов измерений мы неизбежно пользуемся формулами и выводами теории случайных ошибок, и это даёт нам полную возможность вносить большую ясность и определённости при оценке результатов измерений даже в самых сложных случаях.

Общие закономерности, которые мы приписываем случайным ошибкам, очень ясно выявляются, если принять во внимание глубокую аналогию, наблюдаемую между случайными ошибками измерений и ещё одним явлением, которое, как оказывается, также подчиняется только законам случая. Это—стрельба в цель, хотя бы, например, из винтовки. Практика показывает, что если мы производим ряд выстрелов из винтовки, прицеливаясь одинаково тщательно при каждом отдельном выстреле в центр мишени, то пули всегда ложатся не точно в её центр, а на некоторых расстояниях от него (рис. 2), т. е. всегда наблюдается некоторый разброс или, как обычно говорят, некоторое рассеяние точек поражения мишени. Вместе с тем, однако, нетрудно заметить, что точки поражения мишени всегда группируются около её центра, т. е. точки поражения, далёкие от центра мишени, встречаются реже по сравнению с точками, близкими к её центру. Это обстоятельство особенно отчётливо выявляется при большом числе последовательных выстрелов.

Совершенно такую же закономерность мы допускаем и в области случайных ошибок, т. е. мы допускаем, что случайные ошибки, большие по абсолютной величине,

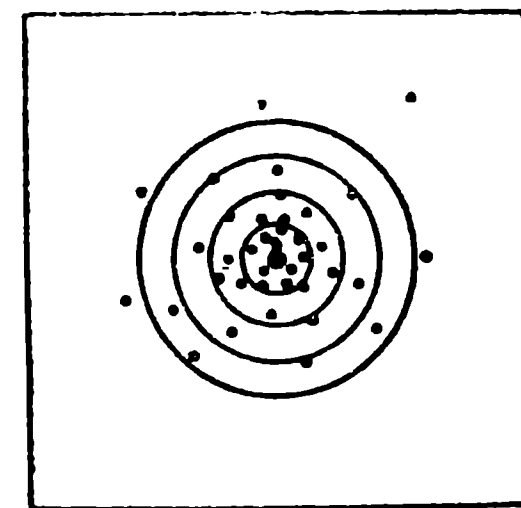


Рис. 2.

должны реже встречаться по сравнению с малыми ошибками; таким образом, случайные ошибки должны группироваться около точного значения измеряемой величины так же, как группируются точки поражения мишени около её центра.

5. Основная формула теории случайных ошибок. Мы можем поставить вопрос о том законе, которому следует распределение случайных ошибок в зависимости от их величины. При этом, предполагая, что случайные ошибки, малые по абсолютной величине, встречаются чаще, чем большие, мы, очевидно, можем сказать, что число, или частота, случайных ошибок является убывающей функцией их величины. Таким образом, эту функцию y , так называемую функцию распределения случайных ошибок, можно на основании сказанного в общем виде представить выражением

$$y = \varphi(x), \quad (13)$$

где x обозначает величину случайных ошибок.

Допустим, что мы произвели n измерений некоторой величины, все одинаковой точности; случайные ошибки этих измерений обозначим через $x_1, x_2, x_3, \dots, x_n$. Мы вправе утверждать, что число случайных ошибок, которые лежат в некоторых пределах между x и $x + dx$, т. е. в интервале dx , должно быть пропорционально:

- 1) общему числу ошибок n ,
- 2) величине взятого интервала dx и
- 3) функции (13) распределения случайных ошибок.

Таким образом, обозначая это число, или частоту ошибок, через dn , можно написать:

$$dn = n\varphi(x)dx, \quad (14)$$

где функция распределения случайных ошибок $\varphi(x)$ остаётся неизвестной.

Эта проблема, т. е. определение вида функции распределения случайных ошибок, была разрешена Гауссом, который создал очень полно и подробно разработанную теорию случайных ошибок, прекрасно отвечающую опыт-

ным данным. Теория Гаусса приводит к так называемому закону нормального распределения случайных ошибок, что даёт возможность исследовать и оценить точность измерений и точность их окончательного результата, а также ближе подойти к установлению точного значения измеряемой величины.

Выводы этой теории основаны на двух общих положениях, которые принято рассматривать как аксиомы. Эти положения, очевидность которых при большом числе измерений одинаковой точности не вызывает сомнений, можно формулировать так:

1) Вероятность (или частота) появления малых случайных ошибок больше вероятности (или частоты) появления больших случайных ошибок. Иначе говоря, вероятность (частота) появления случайных ошибок есть убывающая функция их величины.

2) Вероятность (частота) появления случайных ошибок не зависит от их знака, т. е. случайные ошибки, равные по абсолютной величине, но противоположные по знаку, встречаются одинаково часто.

Первым из этих положений мы уже пользовались; прямым следствием этого положения является уравнение (13), которое теперь на основании второго положения мы можем несколько изменить. Действительно, из второго положения мы видим, что функция $\varphi(x)$ должна быть функцией чётной по отношению к x , т. е. x в правой части уравнения может входить только в чётных степенях, так как перемена знака при x не должна влиять на величину функции $\varphi(x)$. Таким образом, вместо уравнения (13) мы можем написать:

$$y = f(x^2). \quad (15)$$

В соответствии с этим уравнение (14) получает такой вид:

$$dn = n \cdot f(x^2) dx. \quad (16)$$

Из этого выражения мы можем определить вероятность ошибки, равной x , точнее говоря, вероятность того,

что некоторая ошибка лежит в пределах между x и $x + dx$. Действительно, число случаев, благоприятных этому событию, очевидно, равно числу ошибок, лежащих в тех же пределах, т. е. равно dn , а общее число всех возможных случаев равно общему числу случайных ошибок, т. е. равно n . Таким образом, обозначая через dP вероятность того, что некоторая случайная ошибка лежит в пределах между x и $x + dx$, мы можем на основании уравнений (1) и (16) написать:

$$dP = \frac{dn}{n} = f(x^2) dx, \quad (17)$$

где функция распределения случайных ошибок $f(x^2)$ продолжает оставаться попрежнему неизвестной.

Мы определим вид этой функции для рассеяния точек поражения мишени, т. е. найдём закон нормального распределения точек поражения мишени и будем считать его одновременно законом нормального распределения случайных ошибок, пользуясь той полной аналогией между этими явлениями, о которой было сказано выше.

Пусть в центре мишени (рис. 3, а) находится начало прямоугольной системы координат xOy .

Каждая точка поражения мишени, например точка B , будет, очевидно, характеризоваться определёнными координатами. Найдём вероятность того, что точка B лежит внутри бесконечно узкой полоски шириной dx , расположенной на расстоянии x от оси y . Обозначая эту вероятность через P_x , мы можем на основании формулы (17) написать:

$$P_x = f(x^2) dx.$$

Точно так же, обозначая через P_y вероятность того, что точка B лежит в бесконечно узкой полоске шириной dy , расположенной на расстоянии y от оси x , мы можем написать на основании той же формулы:

$$P_y = f(y^2) dy.$$

Вероятность того, что точка B лежит в обеих этих полосках одновременно, иначе говоря, лежит внутри бесконечно малого прямоугольника со сторонами, рав-

ными dx и dy , является вероятностью совместного возникновения обоих указанных событий, т. е. определяется произведением вероятностей P_x и P_y . Таким образом, обозначая эту вероятность через P_1 , находим:

$$P_1 = P_x \cdot P_y = f(x^2) \cdot f(y^2) dx dy. \quad (18)$$

Возьмём вторую систему прямоугольных координат $\xi O \zeta$ с началом также в центре мишени (рис. 3, б) и вычис-

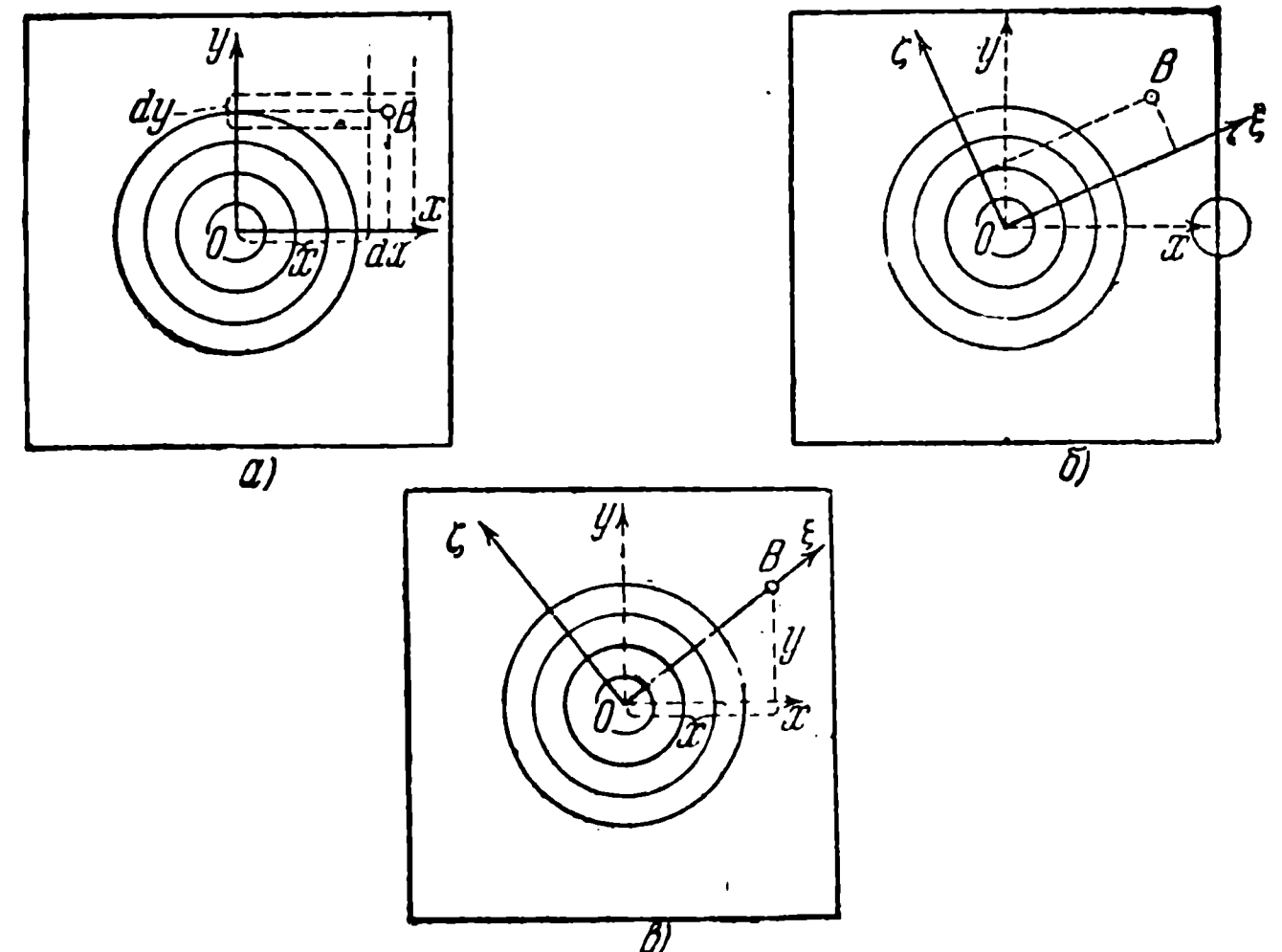


Рис. 3.

лим совершенно тем же приёмом вероятность того, что та же точка B поражения мишени лежит внутри бесконечно малого прямоугольника со сторонами, равными $d\xi$ и $d\zeta$. Мы получим выражение, совершенно аналогичное выражению (18). Обозначая эту вероятность через P_2 , мы можем написать:

$$P_2 = f(\xi^2) \cdot f(\zeta^2) d\xi \cdot d\zeta. \quad (18')$$

Величины $d\xi$ и $d\zeta$ выберем так, чтобы они были равны соответственно величинам dx и dy , т. е. положим, что

$$dx = d\xi \text{ и } dy = d\zeta.$$

При этом условии площади обоих бесконечно малых прямоугольников $(dx \cdot dy)$ и $(d\xi \cdot d\zeta)$ также будут равны. Таким образом имеем:

$$dx \cdot dy = d\xi \cdot d\zeta. \quad (19)$$

В таком случае можно утверждать, что вероятности P_1 и P_2 также должны быть равны, так как вероятность того, что точка поражения мишени B лежит внутри двух равновеликих прямоугольников, не может зависеть от выбора системы координат, иначе говоря, эта вероятность должна остаться одной и той же в любой системе координат. Таким образом, можно положить:

$$f(x^2) \cdot f(y^2) dx \cdot dy = f(\xi^2) \cdot f(\zeta^2) d\xi \cdot d\zeta.$$

Отсюда в соответствии с выражением (19) получаем:

$$f(x^2) \cdot f(y^2) = f(\xi^2) \cdot f(\zeta^2).$$

Расположим теперь систему координат $\xi O \zeta$ так, чтобы одна из её осей, например ось ξ , проходила через точку B (рис. 3, в) и выразим координаты точки B в этой системе через её координаты в системе $x O y$. Из рисунка видно, что

$$\xi^2 = x^2 + y^2 \text{ и } \zeta = 0.$$

Подставляя эти значения ξ и ζ в последнее уравнение, находим:

$$f(x^2) \cdot f(y^2) = f(x^2 + y^2) \cdot f(0). \quad (20)$$

Функция $f(0)$ должна быть постоянной величиной, поэтому можно положить:

$$f(0) = k,$$

где k обозначает некоторую постоянную величину.

В соответствии с этим последнее уравнение принимает такой вид:

$$f(x^2) \cdot f(y^2) = k f(x^2 + y^2). \quad (20')$$

Мы получили так называемое функциональное уравнение, которое даёт возможность определить вид функции f .

Для этого введём вместо переменных x и y новые переменные u и v , полагая:

$$x^2 = u \text{ и } y^2 = v. \quad (21)$$

Получаем:

$$f(u) \cdot f(v) = k f(u + v).$$

Находим частные производные этого выражения по u и по v :

$$\frac{\partial f(u)}{\partial u} \cdot f(v) = k \frac{\partial f(u+v)}{\partial u} \text{ и } \frac{\partial f(v)}{\partial v} \cdot f(u) = k \frac{\partial f(u+v)}{\partial v}.$$

Правые части этих уравнений равны, поэтому приравняем их левые части:

$$\frac{\partial f(u)}{\partial u} \cdot f(v) = \frac{\partial f(v)}{\partial v} \cdot f(u) \text{ или } \frac{\frac{\partial f(u)}{\partial u}}{f(u)} = \frac{\frac{\partial f(v)}{\partial v}}{f(v)}.$$

Обозначим это отношение через b :

$$\frac{\frac{\partial f(u)}{\partial u}}{f(u)} = b.$$

Умножая обе части этого выражения на du , получаем:

$$\frac{\frac{\partial f(u)}{\partial u} \cdot du}{f(u)} = b du.$$

Интегрируем это выражение по u :

$$\ln f(u) = bu + \ln C,$$

где C — постоянная интегриации.

Переходя в последнем выражении от логарифмов к числам, получаем:

$$f(u) = C \cdot e^{bu}.$$

Так как функция f должна быть функцией убывающей, как об этом было сказано выше, то ясно, что величина b должна быть отрицательной; поэтому, полагая:

$$b = -h^2,$$

получаем:

$$f(u) = C e^{-h^2 u}.$$

Вставляя в это выражение вместо u его значение из выражения (21), получаем:

$$f(x^2) = Ce^{-h^2x^2}. \quad (22)$$

В этом выражении величина C остаётся неизвестной. Для того чтобы определить эту величину, воспользуемся уравнением (17). Подставляя в него полученное значение $f(x^2)$, находим:

$$dP = Ce^{-h^2x^2} dx.$$

Это выражение определяет вероятность некоторой ошибки, равной x , т. е. вероятность того, что эта ошибка лежит в пределах между x и $x + dx$. Если мы возьмём интеграл этого выражения в пределах от $-\infty$ до $+\infty$, то мы определим вероятность того, что некоторая ошибка x лежит в пределах от $-\infty$ до $+\infty$. Но эта вероятность, очевидно, является достоверностью, так как несомненно, что все случайные ошибки всех измерений непременно должны лежать в пределах от $-\infty$ до $+\infty$. Вследствие этого мы можем написать:

$$P = \int_{-\infty}^{+\infty} Ce^{-h^2x^2} dx = 1. \quad (23)$$

Этот интеграл преобразуем следующим образом:

$$\int_{-\infty}^{+\infty} Ce^{-h^2x^2} dx = \frac{2C}{h} \int_0^{\infty} e^{-(hx)^2} d(hx). \quad (23')$$

Полагая в последнем интеграле $hx = z$, мы приводим его к виду

$$\int_0^{\infty} e^{-z^2} dz.$$

Но известно, что этот интеграл равен $\frac{\sqrt{\pi}}{2}$, т. е.

$$\int_0^{\infty} e^{-z^2} dz = \frac{\sqrt{\pi}}{2}.$$

На основании этого из уравнений (23) и (23') получаем:

$$\frac{2C}{h} \cdot \frac{\sqrt{\pi}}{2} = 1.$$

Отсюда находим:

$$C = \frac{h}{\sqrt{\pi}}. \quad (24)$$

Вставляя это значение C в уравнение (22) и принимая во внимание уравнение (15), получаем окончательно:

$$y = f(x^2) = \frac{h}{\sqrt{\pi}} e^{-h^2x^2}. \quad (25)$$

Формула (25) даёт выражение закона нормального распределения случайных ошибок. Эта формула, которую

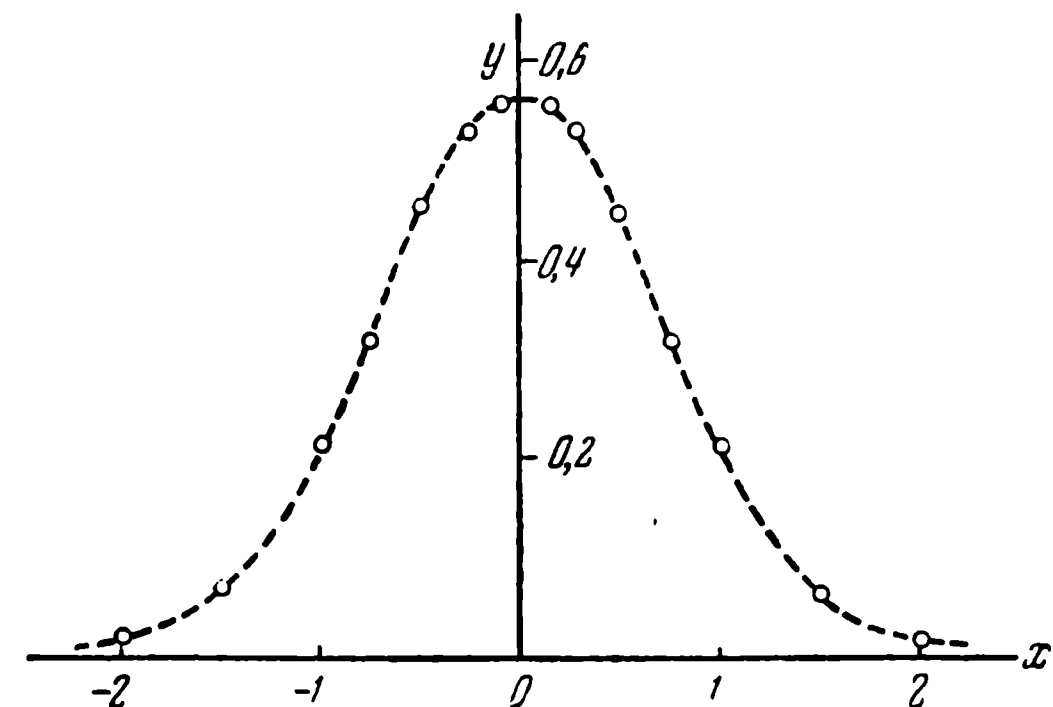


Рис. 4.

называют различно — формулой ошибок, формулой Гаусса, формулой вероятностей, — играет очень большую роль, так как служит основанием всей теории случайных ошибок и теории точности измерений.

Необходимо отметить, что формула (25) содержит только одну постоянную h ; её принято называть мерой точности. Если формулу ошибок изобразить графически (рис. 4), откладывая по оси абсцисс величину случайных ошибок x , а по оси ординат их частоту y , то получается

симметричная по отношению оси ординат колоколообразная кривая, которая в обоих направлениях асимптотически приближается к оси абсцисс. Эта кривая носит различные названия — нормальная кривая вероятностей, кривая ошибок, кривая Гаусса и т. п. Если кривую ошибок вычертить в одном и том же масштабе для различных значений h , то получаются характерные кривые, которые тем выше поднимаются вдоль оси ординат, чем больше оказывается значение h , как это можно видеть

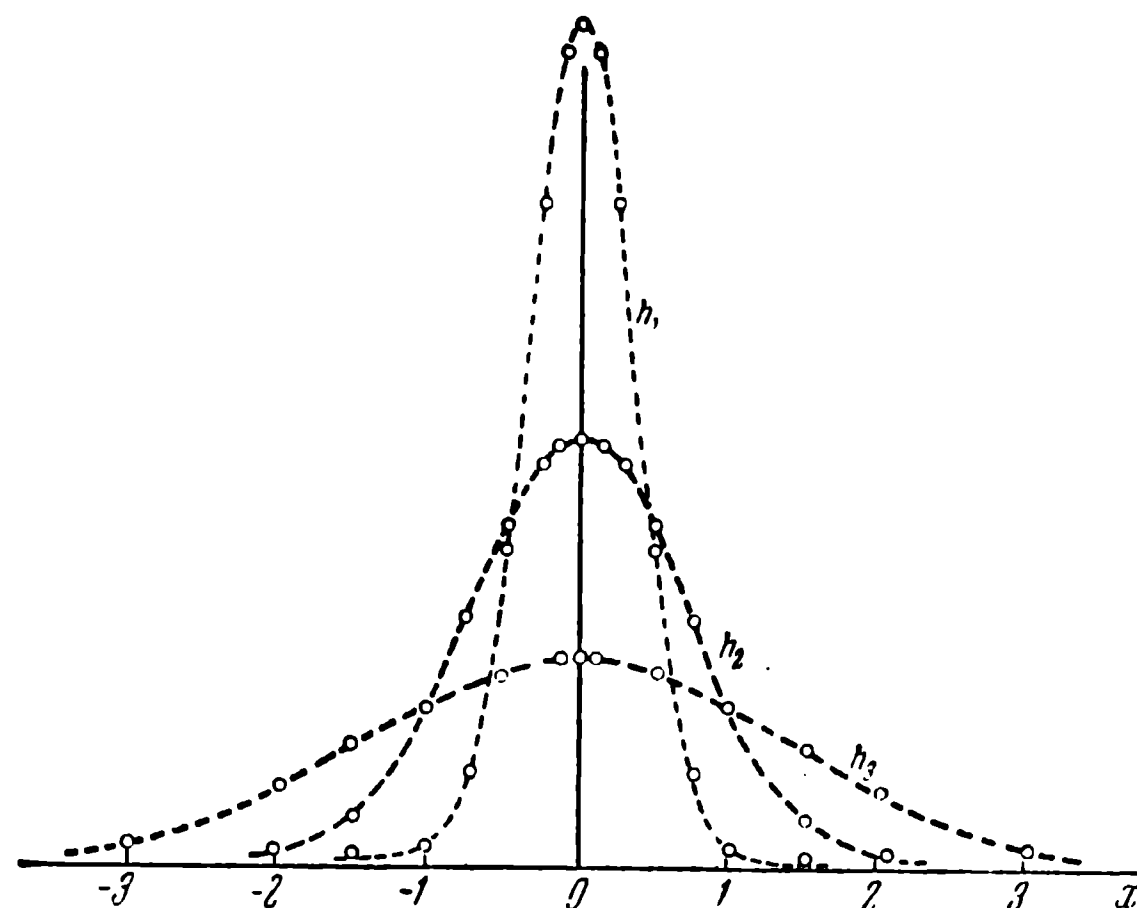


Рис. 5.

из рис. 5. Кривые ошибок для трёх значений h , представленные на рисунке, удовлетворяют условию

$$h_1 > h_2 > h_3.$$

Отсюда становится понятным то название, которое дают величине h , — мера точности. Действительно, чем больше мера точности h , тем круче кривая ошибок падает к оси x . Таким образом, большое значение h говорит о том, что в данной серии измерений резко преобладают случайные ошибки, малые по абсолютной величине, а случайные ошибки, большие по абсолютной вели-

чине, встречаются только в очень небольшом количестве, что является основным признаком измерений высокой точности. Если вновь воспользоваться аналогией со стрельбой в цель, то мы вправе сказать, что при большом значении меры точности h очень большое число точек поражения мишени располагается вблизи её центра и лишь немногие оказываются вдали от него, — такой стрельбе или такому виду оружия мы должны дать высокую оценку.

6. Интеграл вероятностей и его вычисление. Значение y или $f(x^2)$ из уравнения (25) представим в уравнение (17). Находим:

$$dP = \frac{h}{\sqrt{\pi}} e^{-h^2 x^2} dx. \quad (26)$$

Это выражение определяет вероятность того, что некоторая случайная ошибка x лежит в пределах бесконечно

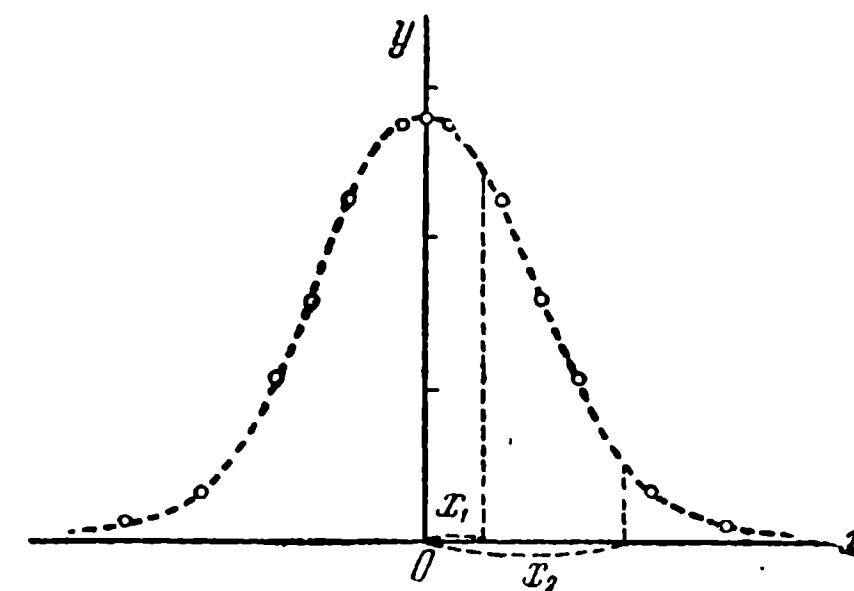


Рис. 6.

малого интервала dx . Возьмём интеграл этого выражения в некоторых пределах от x_1 до x_2 :

$$P = \frac{h}{\sqrt{\pi}} \int_{x_1}^{x_2} e^{-h^2 x^2} dx. \quad (27)$$

Интеграл в правой части выражения (27) изображается, как это показано на рис. 6, площадью, ограниченной кривой ошибок, двумя ординатами, отвечающими абсциссам

x_1 и x_2 , и отрезком оси абсцисс между этими ординатами. Выражение (27) определяет вероятность того, что некоторая случайная ошибка лежит в пределах конечного интервала между x_1 и x_2 . Это выражение, которое носит название интеграла вероятностей, обыкновенно пишут в несколько иной форме.

В выражении (27) оба предела интеграла лежат по одну сторону от начала координат, но можно предположить, что пределами интеграла взяты два значения x , противоположные по знаку, но равные по абсолютной

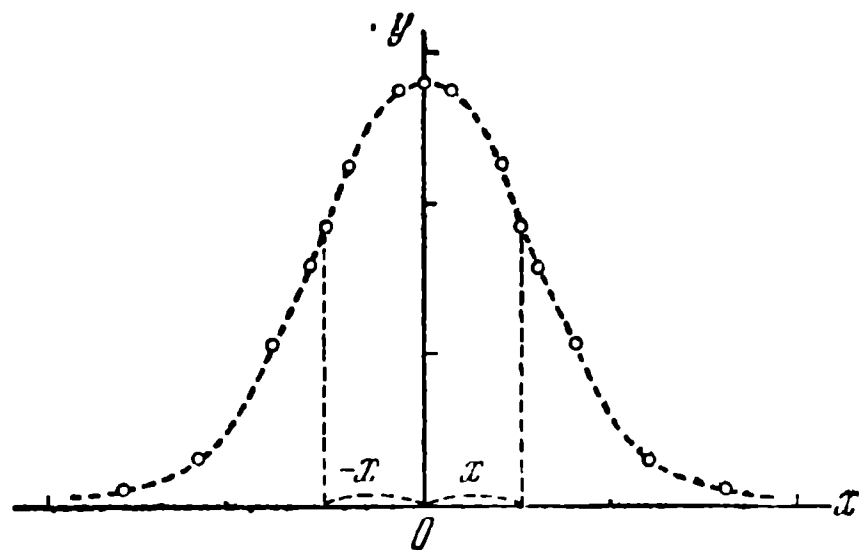


Рис. 7.

величине, как это показано на рис. 7. В таком случае, обозначая эти пределы через $-x$ и $+x$, имеем:

$$P = \frac{h}{\sqrt{\pi}} \int_{-x}^{+x} e^{-h^2 x^2} dx. \quad (28)$$

Интеграл в правой части этого выражения можно, очевидно, представить как сумму двух интегралов с пределами от $-x$ до 0 и от 0 до $+x$, т. е. можно написать:

$$\int_{-x}^{+x} e^{-h^2 x^2} dx = \int_{-x}^0 e^{-h^2 x^2} dx + \int_0^{+x} e^{-h^2 x^2} dx.$$

Отсюда, принимая во внимание, что оба интеграла в правой части этого выражения равны между собой,

находим:

$$\int_{-x}^{+x} e^{-h^2 x^2} dx = 2 \int_0^x e^{-h^2 x^2} dx.$$

Вследствие этого выражение (28) можно представить в таком виде:

$$P = \frac{2h}{\sqrt{\pi}} \int_0^x e^{-h^2 x^2} dx.$$

В этом выражении h вводим под знак интеграла в виде множителя при x . Находим:

$$P = \frac{2}{\sqrt{\pi}} \int_0^x e^{-h^2 x^2} d(hx). \quad (29)$$

Окончательное выражение интеграла вероятностей:

$$P = \frac{2}{\sqrt{\pi}} \int_0^z e^{-z^2} dz, \quad (30)$$

где

$$z = hx.$$

Вычислить интеграл вероятностей обычными приёмами интегрального исчисления не представляется возможным, поэтому для его определения приходится пользоваться методами приближённых вычислений определённых интегралов. Для этой цели в данном случае можно разложить подинтегральную функцию e^{-z^2} в бесконечный ряд:

$$e^{-z^2} = 1 - z^2 + \frac{z^4}{2!} - \frac{z^6}{3!} + \frac{z^8}{4!} - \dots$$

Подставляя это значение e^{-z^2} в правую часть выражения (30), получаем:

$$P = \frac{2}{\sqrt{\pi}} \int_0^z \left(1 - z^2 + \frac{z^4}{2!} - \frac{z^6}{3!} + \frac{z^8}{4!} - \frac{z^{10}}{5!} + \dots \right) dz.$$

Выполняя интегрирование, находим:

$$P = \frac{2}{\sqrt{\pi}} \left(z - \frac{z^3}{3} + \frac{z^5}{10} - \frac{z^7}{42} + \frac{z^9}{216} - \frac{z^{11}}{1320} + \dots \right).$$

Сходимость ряда в правой части этого выражения уменьшается при возрастании значений z , в чём нетрудно непосредственно убедиться. Так для трёх значений z :

$$z = 0,1; \quad z = 0,5; \quad z = 1,0,$$

получаем:

1. При $z = 0,1$

$$P = \frac{2}{\sqrt{\pi}} \left(0,1 - \frac{0,001}{3} + \frac{0,00001}{10} - \dots \right)$$

или

$$P = \frac{2}{\sqrt{\pi}} (0,1 - 0,00033 + 0,000001 - \dots) = 0,1125.$$

Мы видим, что для вычисления P с точностью, например, до четвёртого десятичного знака достаточно ограничиться двумя первыми членами ряда.

2. При $z = 0,5$

$$P = \frac{2}{\sqrt{\pi}} \left(0,5 - \frac{1}{24} + \frac{1}{320} - \frac{1}{5376} + \dots \right)$$

или

$$P = \frac{2}{\sqrt{\pi}} (0,5 - 0,0417 + 0,0031 - 0,0002 + 0,000001 - \dots) = 0,5205.$$

В этом случае для вычисления P с той же точностью, т. е. до четвёртого десятичного знака, необходимо взять четыре или даже пять первых членов ряда.

3. При $z = 1,0$

$$P = \frac{2}{\sqrt{\pi}} \left(1,0 - \frac{1}{3} + \frac{1}{10} - \frac{1}{42} + \frac{1}{216} - \frac{1}{1320} + \dots \right)$$

или

$$P = \frac{2}{\sqrt{\pi}} (1,0 - 0,3333 + 0,1 - 0,0238 + 0,0046 - 0,00076 + 0,0001 - 0,000013 + \dots) = 0,8427.$$

При этом значении z для вычисления P с прежней точностью необходимо взять восемь первых членов ряда.

При дальнейшем увеличении значений z этот ряд для вычисления P оказывается крайне неудобным. В таких случаях пользуются разложением интеграла вероятностей в ряд иного вида, который даёт возможность более быстро вычислить значения P с достаточной степенью точности. Этих вопросов мы рассматривать не будем.

Значения интеграла вероятностей вычислены с большой точностью для различных значений предела интегрирования z , при постепенном изменении их на очень небольшую величину, например для значений z , отличающихся на 0,01 или даже 0,001. Эти значения интеграла вероятностей собраны в особые таблицы (стр. 375), подобные, например, таблицам тригонометрических функций. Таковыми таблицами очень часто приходится пользоваться как при обработке результатов измерений, так и в других практических вопросах, основанных на теории случайных ошибок, например в вопросах пристрелки оружия, в вопросах математической статистики и т. п. Некоторые значения интеграла вероятностей приведены в таблице II.

Таблица II
Значения интеграла вероятностей

z	P	z	P
0,0	0	1,5	0,9661
0,1	0,1125	2,0	0,9953
0,5	0,5205	3,0	0,999978
1,0	0,8427	4,0	0,999999985

Из этой таблицы мы видим, что при увеличении z значения интеграла вероятностей также возрастают, стремясь к пределу, равному единице. Это возрастание значений P вполне понятно, так как при увеличении z , т. е. при расширении предела интегрирования, вероятность того, что некоторая случайная положительная ошибка

будет лежать в этих пределах, очевидно, должна возра-
стать, стремясь к достоверности при бесконечном увели-
чении верхнего предела интеграции.

Теория случайных ошибок, которой мы постоянно
пользуемся при экспериментальных работах, неоднократно
подвергалась весьма серьёзной и ответственной проверке
при обработке, например, результатов очень обширных
астрономических и геодезических измерений; из таких
испытаний теория ошибок всегда выходила с честью.
В таких обширных работах число отдельных измерений
обычно определялось несколькими сотнями, часто оказы-
валось больше тысячи. Обработывая этот громадный число-
вой материал, определяли непосредственным подсчётом
число ошибок различной величины, от наименьшей до наи-
большей, наблюдавшейся при измерениях. Если эти
числа сопоставлять с теми, к которым приводит теория слу-
чайных ошибок, то всегда наблюдалось практически пол-
ное совпадение теоретических чисел с числами, получен-
ными непосредственным подсчётом.

Необходимо ещё указать, что для проверки теории
случайных ошибок часто применялись различные прибо-
ры, в которых искусственно создаётся очень большое
число однородных явлений и они протекают при таких
условиях, что могут подчиняться только законам случая.
Такие условия мы можем создать различными приёмами.
Так, один из приборов, которые можно применять для
проверки теории случайных ошибок, состоит из гладкой
широкой доски, установленной наклонно, в которой укреп-
лено очень большое число штифтиков или гвоздиков.
В верхнем конце доски, против её середины, установлена
воронка, через отверстие которой на доску направляется
непрерывной струйкой очень большое число мелких шариков
(мелкая дробь). Скатываясь по доске, шарики уда-
ряются о штифтики и испытывают взаимные удары, так
что направление их движения непрерывно изменяется.
Весь процесс протекает под действием ряда факторов, не
поддающихся точному учёту, таких, например, как на-
правление движения до столкновения, число ударов, их
сила, их направление и т. д. В результате струйка шариков
испытывает рассеяние, и на некотором расстоянии от

конца воронки вдоль доски начинает двигаться поток
шариков, постепенно расширяющийся при их дальнейшем
движении. Как показывают наблюдения, среднее число
шариков в потоке, рассчитанное на единицу площади, или
средняя плотность потока шариков в различных точках
его поперечного сечения оказывается различной. Плот-
ность оказывается наибольшей в середине потока, против
отверстия воронки, а в обе
стороны от середины по-
тока его плотность сим-
метрично уменьшается. Это
изменение плотности по
направлению к краям по-
тока хорошо отвечает осно-
вному закону теории слу-
чайных ошибок.

В другом приборе того
же типа из отверстия во-
ронки *B* (рис. 8) вытекает
непрерывная струйка мел-
кого отсеянного песка.
Струйка песка, частично
рассеянная после выхода
из воронки вследствие тре-
ния и взаимных столкно-
вений, встречает затем две
или несколько металличе-
ских сеток *СС*, установ-
ленных горизонтально на
небольшом расстоянии
друг от друга. Отверстия
в сетках имеют достаточ-
но большую величину, так
что отдельные песчинки
свободно проходят через
них, но много песчинок
ударяется о проволочки
сеток. Направление дви-
жения этих песчинок
резко изменяется, что
вызывает большое число
взаимных столкновений.
В результате поток
песчинок обнаруживает
сильное рассеяние и на
некотором расстоянии
от сеток становится
достаточно широким.
Если внизу поставить
длинный ящик *M* с
прозрачными стенками,

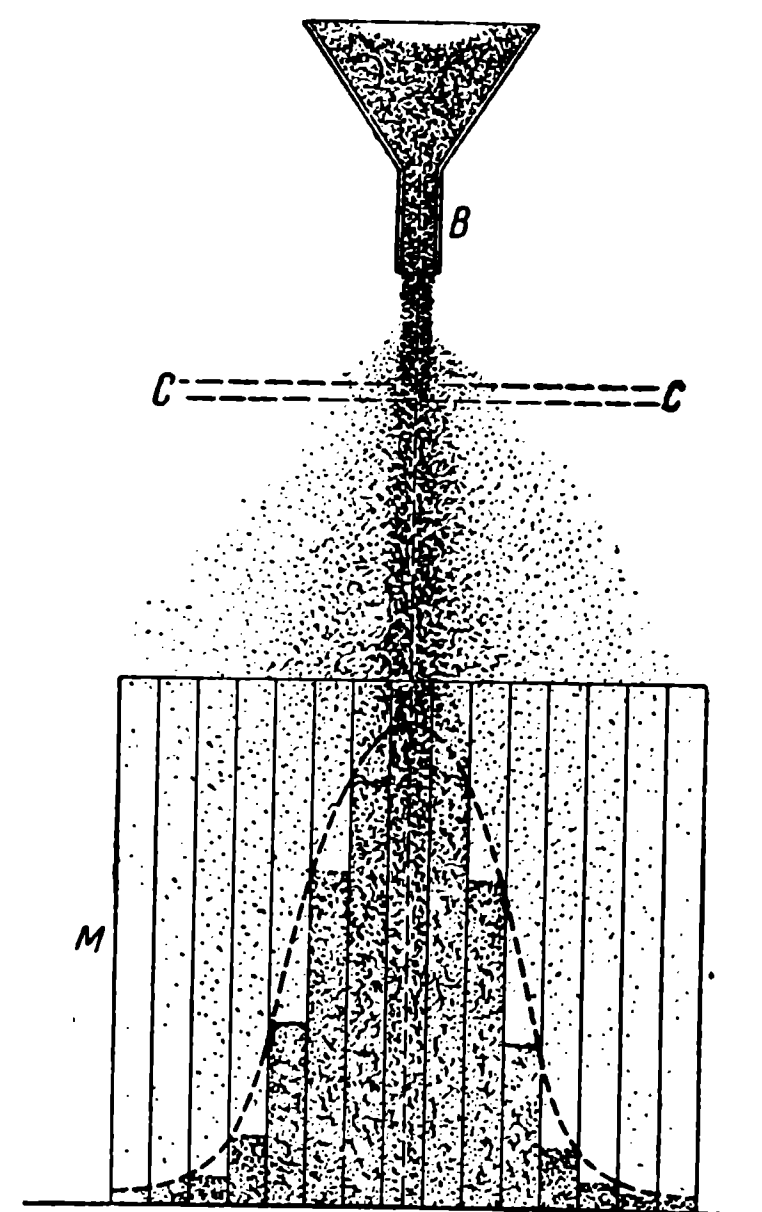


Рис. 8.

разделённый поперечными перегородками на большое число секций, то в различных секциях начинает собираться различное количество песка. В средней секции, которая расположена под отверстием воронки, песок собирается в наибольшем количестве, и песчаный столбик здесь быстро растёт. В остальных секциях песка оказывается тем меньше, чем дальше они расположены от средней секции, причём уменьшение количества песка в обе стороны от средней секции имеет вполне симметричный характер. Получается ясно выраженная картина ступенчатого расположения песка в секциях, которая хорошо отвечает кривой ошибок. На этом приборе можно проследить изменение кривой ошибок при изменении меры точности h . Эта последняя величина в данном случае зависит от рассеяния потока песчинок, которое в этом приборе можно изменять, например, в зависимости от числа сеток CC . Увеличивая их число, мы увеличиваем рассеяние песчинок, что вызывает уменьшение меры точности h . Непосредственными наблюдениями нетрудно убедиться в том, что это действительно имеет место: при увеличении числа сеток получаются кривые ошибок более отлогие, с максимумом, выраженным менее резко.

ГЛАВА V

ПОКАЗАТЕЛИ ТОЧНОСТИ ИЗМЕРЕНИЙ

1. Постулат среднего арифметического. Закон нормального распределения случайных ошибок даёт возможность разъяснить те основания, в силу которых принято считать, что окончательное значение измеряемой величины определяется средним арифметическим из всех отдельных измерений. Такое обоснование постулата среднего арифметического можно сделать на основании следующих соображений.

Допустим, что мы произвели n измерений некоторой величины. Обозначим результаты этих измерений через $a_1, a_2, a_3, \dots, a_n$, а случайные ошибки измерений соответственно через $x_1, x_2, x_3, \dots, x_n$. Мы предполагаем, что все измерения выполнены на одних и тех же приборах с одинаковой точностью. Неизвестное нам точное значение измеряемой величины обозначим через X . В таком случае величины случайных ошибок измерений могут быть представлены выражениями (стр. 21):

$$\begin{aligned} x_1 &= X - a_1, \\ x_2 &= X - a_2, \\ x_3 &= X - a_3, \\ &\dots \dots \dots, \\ &\dots \dots \dots, \\ x_n &= X - a_n, \end{aligned} \tag{1}$$

где величина ошибок $x_1, x_2, x_3, \dots, x_n$, а также величина X остаются неизвестными.

Вероятности ошибок $x_1, x_2, x_3, \dots, x_n$ можно определить на основании формулы (26) (стр. 151), причём мера точности h для всех измерений будет одной и той же, так как мы предполагаем, что все измерения выполнены с одинаковой точностью. Таким образом, обозначая эти вероятности соответственно через $dP_1, dP_2, dP_3, \dots, dP_n$, мы можем написать:

$$\begin{aligned} dP_1 &= \frac{h}{\sqrt{\pi}} e^{-h^2 x_1^2} dx_1, \\ dP_2 &= \frac{h}{\sqrt{\pi}} e^{-h^2 x_2^2} dx_2, \\ dP_3 &= \frac{h}{\sqrt{\pi}} e^{-h^2 x_3^2} dx_3, \\ &\dots \dots \dots, \\ dP_n &= \frac{h}{\sqrt{\pi}} e^{-h^2 x_n^2} dx_n. \end{aligned} \quad (2)$$

Вероятность появления всех случайных ошибок $x_1, x_2, x_3, \dots, x_n$ следует рассматривать как вероятность сложного события (стр. 136), т. е. совместного появления всех событий, вероятности которых определяются выражениями (2). Вследствие этого, обозначая вероятность появления всех случайных ошибок (1) через P , мы можем на основании теоремы умножения вероятностей написать:

$$P = \frac{h}{\sqrt{\pi}} e^{-h^2 x_1^2} dx_1 \cdot \frac{h}{\sqrt{\pi}} e^{-h^2 x_2^2} dx_2 \times \\ \times \frac{h}{\sqrt{\pi}} e^{-h^2 x_3^2} dx_3 \dots \frac{h}{\sqrt{\pi}} e^{-h^2 x_n^2} dx_n.$$

Отсюда получаем:

$$P = \left(\frac{h}{\sqrt{\pi}} \right)^n e^{-h^2(x_1^2 + x_2^2 + x_3^2 \dots + x_n^2)} dx_1 dx_2 dx_3 \dots dx_n. \quad (3)$$

Так как ошибки $x_1, x_2, x_3, \dots, x_n$ остаются неизвестными, то точного значения измеряемой величины X мы не можем найти. Вследствие этого мы принуждены ограничиться задачей, практически возможной: определить на основании результатов наших измерений приближённое,

но наиболее близкое к точному значение измеряемой величины, т. е. наилучшее значение измеряемой величины. Однако при этом мы должны иметь в виду, что с точки зрения применяемых нами формул наилучшее значение измеряемой величины мы можем рассматривать только как её наиболее вероятное значение. Это наилучшее или наиболее вероятное значение измеряемой величины обозначим через A и найдём выражение, которое определяет A как функцию результатов отдельных измерений, т. е. величин $a_1, a_2, a_3, \dots, a_n$.

Очевидно, что наиболее вероятное значение измеряемой величины A должно соответствовать максимуму вероятности P в выражении (3). Но максимум P , как это нетрудно видеть, отвечает минимуму показателя степени при e в правой части этого выражения, а минимум этого показателя при определённом значении h отвечает минимуму суммы квадратов ошибок измерений x . Таким образом, обозначая сумму квадратов x через z , т. е. полагая:

$$z = x_1^2 + x_2^2 + x_3^2 + \dots + x_n^2 = \sum x^2, \quad (4)$$

мы вправе утверждать, что наиболее вероятным или наилучшим значением измеряемой величины должно быть такое значение A , для которого сумма квадратов ошибок будет иметь наименьшее значение.

Для того чтобы определить соответствующее значение A , мы подставляем эту величину в уравнения (1) вместо неизвестной нам величины X и полученные отсюда значения ошибок $x_1, x_2, x_3, \dots, x_n$ вставляем в выражение (4). Находим:

$$z = (A - a_1)^2 + (A - a_2)^2 + (A - a_3)^2 + \dots + (A - a_n)^2.$$

Отсюда определяем то значение A , при котором это выражение имеет минимум. Для этого находим первую производную z по A и приравняем её нулю:

$$\begin{aligned} \frac{dz}{dA} &= 2(A - a_1) + 2(A - a_2) + 2(A - a_3) + \dots \\ &\dots + 2(A - a_n) = 0. \end{aligned}$$

Из этого выражения после небольших преобразований находим:

$$nA - (a_1 + a_2 + a_3 + \dots + a_n) = 0$$

или

$$A = \frac{a_1 + a_2 + a_3 + \dots + a_n}{n} = \frac{\sum a}{n}. \quad (5)$$

Мы получили среднее арифметическое из всех измерений. Таким образом, мы вправе сказать, что:

При измерениях одинаковой точности наиболее вероятным или наилучшим значением измеряемой величины является среднее арифметическое из результатов, полученных при всех измерениях.

Это положение или так называемый постулат среднего арифметического и применяется при всех вычислениях окончательного результата измерений. Теперь мы видим, что постулат среднего арифметического является прямым следствием теории случайных ошибок. При этом никогда не следует упускать из вида, что среднее арифметическое, как показывает наш вывод, является только наиболее вероятным значением измеряемой величины. Следует также обратить внимание на выражение (4). Это выражение определяет общие условия, которым должно удовлетворять значение A : надо найти такое значение A , для которого сумма квадратов случайных ошибок имеет наименьшее значение. Это положение лежит в основе так называемого принципа наименьших квадратов, который находит очень широкое применение, например, применяется во всех случаях, когда приходится решать систему уравнений, в которой число неизвестных меньше числа уравнений.

2. Мера точности для точных значений ошибок и для отклонений от среднего арифметического. При выводе формулы (5) мы заменили в уравнениях (1) точное значение измеряемой величины X величиной A , которая оказалась равной среднему арифметическому результатов всех измерений. Таким образом, в правых частях урав-

нений (1) появились разности между средним арифметическим результатов измерений и каждым из них в отдельности. Эти разности принято называть отклонениями отдельных измерений от их среднего арифметического; мы будем обозначать их в дальнейшем через ϵ с соответствующим индексом. При выводе постулата среднего арифметического мы допустили, что величины ϵ равны точным значениям ошибок измерений x . Но такое предположение, строго говоря, мы должны считать не вполне точным: величины ϵ и x являются, несомненно, очень близкими, но они не могут быть в точности равными. Кроме того, необходимо иметь в виду, что все основные положения теории случайных ошибок мы вывели для величин x , т. е. для точных значений ошибок измерений, и было бы рискованным распространять эти положения без особых доказательств также и на величины ϵ , т. е. на отклонения результатов измерений от их среднего арифметического. Так как на практике мы должны заменять неизвестные нам величины x величинами ϵ , которые всегда можно вычислить, то, очевидно, необходимо выяснить, не следует ли внести в наши выводы какие-либо изменения при переходе от точных значений ошибок измерений x к их отклонениям от среднего арифметического ϵ .

В теории случайных ошибок эти вопросы были подробно исследованы, и те затруднения, которые здесь возникают, оказалось возможным устранить. Мы познакомимся с результатами этих исследований только в основных чертах, принимая некоторые из положений без доказательств.

Допустим попрежнему, что мы произвели n измерений некоторой величины, все одинаковой точности. Введём прежние обозначения, т. е. обозначим результаты отдельных измерений через $a_1, a_2, a_3, \dots, a_n$, их среднее арифметическое или наиболее вероятное значение измеряемой величины через A , а отклонения результатов отдельных измерений от их среднего арифметического обозначим соответственно через $\epsilon_1, \epsilon_2, \epsilon_3, \dots, \epsilon_n$; наконец, через X обозначим точное значение измеряемой величины, нам неизвестное, и через $x_1, x_2, x_3, \dots, x_n$ точные значения ошибок всех измерений, также нам неизвестные.

Мы можем попрежнему написать ряд равенств (1) для точных значений ошибок измерений x и совершенно аналогичный ряд равенств для отклонений ε от среднего арифметического:

$$\begin{aligned} x_1 &= X - a_1, & \varepsilon_1 &= A - a_1, \\ x_2 &= X - a_2, & \varepsilon_2 &= A - a_2, \\ x_3 &= X - a_3, & \varepsilon_3 &= A - a_3, \\ &\dots\dots\dots, & & \\ &\dots\dots\dots, & & \\ x_n &= X - a_n; & \varepsilon_n &= A - a_n. \end{aligned} \quad (6)$$

Складывая все равенства в том и другом ряду в отдельности и вводя для сокращения знак суммы $\sum$, находим:

$$\begin{aligned} \sum x &= nX - \sum a, \\ \sum \varepsilon &= nA - \sum a. \end{aligned}$$

Отсюда находим:

$$\sum x - \sum \varepsilon = nX - nA.$$

Величина $\sum \varepsilon$ представляет собой алгебраическую сумму отклонений результатов всех измерений от их среднего арифметического. Нетрудно доказать для любого ряда величин, что алгебраическая сумма отклонений этих величин от их среднего арифметического всегда равна нулю; это является основным свойством среднего арифметического. Таким образом, мы можем положить:

$$\sum \varepsilon = 0.$$

Вследствие этого последнее равенство получает такой вид:

$$\sum x = nX - nA.$$

Определяя из этого выражения величину A , находим:

$$A = X - \frac{\sum x}{n}. \quad (7)$$

Это значение A подставляем в правый ряд равенств (6) и, принимая во внимание левый ряд тех же равенств,

получаем после небольших алгебраических преобразований ряд таких выражений:

$$\begin{aligned} \varepsilon_1 &= \frac{n-1}{n} x_1 - \frac{x_2}{n} - \frac{x_3}{n} - \dots - \frac{x_n}{n}, \\ \varepsilon_2 &= -\frac{x_1}{n} + \frac{n-1}{n} x_2 - \frac{x_3}{n} - \dots - \frac{x_n}{n}, \\ \varepsilon_3 &= -\frac{x_1}{n} - \frac{x_2}{n} + \frac{n-1}{n} x_3 - \dots - \frac{x_n}{n}, \\ &\dots\dots\dots, \\ \varepsilon_n &= -\frac{x_1}{n} - \frac{x_2}{n} - \frac{x_3}{n} - \dots + \frac{n-1}{n} x_n. \end{aligned} \quad (8)$$

Эти выражения показывают, что отклонения результатов измерений от среднего арифметического являются линейными функциями точных значений ошибок измерений. Отсюда также следует, что справедливо и обратное заключение, т. е. что точные значения ошибок x также являются линейными функциями отклонений ε .

Рассматривая основные вопросы теории ошибок, оказалось возможным доказать, что если закон нормального распределения ошибок и следствия, из него вытекающие, применимы к некоторому ряду случайных величин, удовлетворяющих условию независимости, то тот же закон и его следствия оказываются применимыми ко всяким линейным функциям этих величин. При этом оказывается, что если линейная функция E ряда независимых случайных величин $a_1, a_2, a_3, \dots, a_n$, следующих закону нормального распределения, определяется выражением

$$E = b_1 a_1 + b_2 a_2 + b_3 a_3 + \dots + b_n a_n,$$

где $b_1, b_2, b_3, \dots, b_n$ являются некоторыми постоянными коэффициентами, то мера точности H функции E определяется формулой

$$\frac{1}{H^2} = \frac{b_1^2 + b_2^2 + b_3^2 + \dots + b_n^2}{h^2}, \quad (9)$$

где h обозначает меру точности величин $a_1, a_2, a_3, \dots, a_n$.

Эти положения мы не будем разбирать подробно и пока примем их без доказательств (ср. стр. 184 и след.).

Таким образом, мы вправе утверждать, что, поскольку закон нормального распределения применим к величинам x , т. е. к точным значениям ошибок измерений, то тот же закон должен быть в той же мере применим и к величинам ϵ , т. е. к отклонениям результатов измерений от их среднего арифметического.

Выражение (9) даёт нам возможность вычислить меру точности отклонений ϵ , если дана мера точности точных значений ошибок x , и наоборот. Обозначим через H меру точности отклонений ϵ и через h меру точности точных значений ошибок x . На основании выражения (9), вводя в него вместо коэффициентов $b_1, b_2, b_3, \dots, b_n$ коэффициенты при $x_1, x_2, x_3, \dots, x_n$, взятые из первого уравнения (8), мы можем написать:

$$\frac{1}{H^2} = \frac{\left(\frac{n-1}{n}\right)^2 + \left(\frac{1}{n}\right)^2 + \left(\frac{1}{n}\right)^2 + \dots + \left(\frac{1}{n}\right)^2}{h^2},$$

т. е.

$$\frac{1}{H^2} = \frac{1}{h^2} \left[\left(\frac{n-1}{n}\right)^2 + \frac{n-1}{n^2} \right]$$

или

$$\frac{1}{H^2} = \frac{1}{h^2} \left(\frac{n^2 - n}{n^2} \right).$$

Отсюда находим:

$$H = h \sqrt{\frac{n}{n-1}}. \quad (10)$$

Из этого выражения мы можем сделать ряд выводов:

1) Во-первых, мы видим, что мера точности H отклонений ϵ результатов измерений от их среднего арифметического и мера точности h точных значений ошибок x оказываются различными, причём всегда

$$H > h, \quad (11)$$

так как величина под знаком корня в выражении (10), очевидно, всегда больше единицы:

$$\frac{n}{n-1} > 1.$$

Таким образом, производя все вычисления по необходимости с величинами ϵ , мы получаем меру точности, бо́льшую той, которая отвечает точным значениям ошибок измерений x , т. е. меру точности больше действительной, иначе говоря, мы как бы преувеличиваем точность наших измерений.

2) Во-вторых, мера точности H отклонения ϵ является функцией меры точности h точных значений ошибок x и числа измерений n ; таким образом, при одной и той же мере точности h мера точности отклонений ϵ может быть различной в зависимости от числа измерений.

3) Наконец, при увеличении числа измерений n , величина, стоящая под знаком корня, постепенно приближается к единице. Таким образом, при увеличении числа измерений мера точности H отклонений ϵ постепенно приближается к мере точности h точных значений ошибок x , но делается ей равной только при бесконечно большом числе измерений. Практически та и другая мера точности становятся равными при достаточно большом числе измерений.

В соответствии с этими выводами очевидно, что основное уравнение вероятностей (стр. 149) следует писать различно для точных значений ошибок x и для отклонений ϵ , а именно, обозначая функцию y для точных значений ошибок через y_x и для отклонений через y_ϵ , мы можем написать:

$$y_x = \frac{h}{\sqrt{\pi}} e^{-h^2 x^2}, \quad (12)$$

$$y_\epsilon = \frac{H}{\sqrt{\pi}} e^{-H^2 \epsilon^2}, \quad (12')$$

где H определяется выражением (10).

Для того чтобы несколько лучше познакомиться с последними выводами, рассмотрим числовой пример.

Примем, например, условно, что при производстве некоторых измерений мера точности ошибок оказалась равной двум, т. е. $h=2$. Допустим, далее, что мы произвели вначале пять измерений, затем увеличили их число до десяти, затем до двадцати и, наконец, довели число измерений до ста. Если, пользуясь формулой (10), мы

вычисляем значение меры точности H отклонений ϵ для всех этих случаев, то получаются результаты, которые приведены в таблице III.

Из этих результатов мы видим, что при небольшом числе измерений мера точности отклонений ϵ заметно

Таблица III

n	$h=2$	H
	$\sqrt{\frac{n}{n-1}}$	
5	1,12	2,24
10	1,05	2,10
20	1,02 (5)	2,05
100	1,005	2,01

отличается от меры точности ошибок x ; так, при пяти измерениях это различие достигает 12%. При увеличении числа измерений это различие уменьшается, и при числе измерений, равном ста, меру точности отклонений ϵ и меру точности ошибок x можно считать практически равными, так как различие между ними в данном случае составляет всего лишь 0,5%.

Если для случая пяти измерений графически изобразить, пользуясь прежней системой координат, кривые

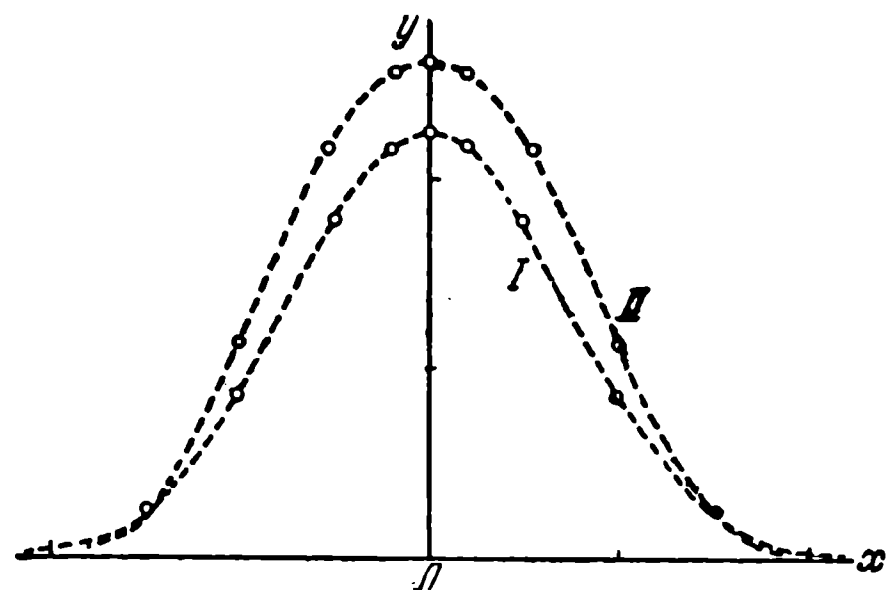


Рис. 9.

вероятностей для точных значений ошибок, т. е. функцию (12), и для отклонений, т. е. функцию (12'), то между обеими кривыми получается очень заметное расхождение, как это ясно видно из рис. 9. Кривая I отвечает величине h , т. е. мере точности ошибок x ; кривая II отвечает вели-

чине H , т. е. той кривой ошибок, к которой мы приходим, оперируя с отклонениями ϵ результатов измерений от их среднего арифметического.

Установим, далее, зависимость между суммой квадратов отклонений ϵ от среднего арифметического, т. е. величиной $\sum \epsilon^2$, и суммой квадратов ошибок x , т. е. величиной $\sum x^2$, по отношению которой мы имели очень важный вывод (стр. 161).

Для этого вновь обратимся к уравнениям (6). Исключая из них величины $a_1, a_2, a_3, \dots, a_n$, приводим эти уравнения к такому виду:

$$\begin{aligned}\epsilon_1 &= x_1 + (A - X), \\ \epsilon_2 &= x_2 + (A - X), \\ \epsilon_3 &= x_3 + (A - X), \\ &\dots \dots \dots \\ \epsilon_n &= x_n + (A - X).\end{aligned}$$

Определив из выражения (7) значение разности $(A - X)$, вставляем его в правые части последних уравнений. Находим:

$$\begin{aligned}\epsilon_1 &= x_1 - \frac{\sum x}{n}, \\ \epsilon_2 &= x_2 - \frac{\sum x}{n}, \\ \epsilon_3 &= x_3 - \frac{\sum x}{n}, \\ &\dots \dots \dots \\ \epsilon_n &= x_n - \frac{\sum x}{n}.\end{aligned} \tag{13}$$

Эти уравнения возводим в квадрат и складываем. Первая операция, т. е. возведение в квадрат уравнений (13), является достаточно кропотливой, если её производить вполне точно. Вследствие этого мы произведём эту операцию приближённо, отбрасывая при возведении в квадрат правых частей уравнений (13) все удвоенные произведения членов, т. е. считая, что алгебраическую сумму произведений первых степеней случайных ошибок $x_1, x_2,$

$x_3, \dots, x_n$, взятых попарно, можно принять равной нулю. Этот вывод можно сделать на основании второго основного положения теории случайных ошибок (стр. 143), согласно которому случайные ошибки, равные по абсолютной величине, но противоположные по знаку, встречаются одинаково часто. Отсюда следует, что и их попарные произведения должны давать нам ряд таких чисел, среди которых числа, равные по абсолютной величине, но противоположные по знаку, должны встречаться также одинаково часто, и сумма всех этих чисел может быть принята равной нулю. При достаточно большом числе измерений одинаковой точности это положение должно давать достаточно точные результаты.

В таком случае в результате возведения в квадрат уравнений (13) мы получаем такой ряд равенств:

$$\varepsilon_1^2 = x_1^2 - 2 \frac{x_1^2}{n} + \frac{\sum x^2}{n^2},$$

$$\varepsilon_2^2 = x_2^2 - 2 \frac{x_2^2}{n} + \frac{\sum x^2}{n^2},$$

$$\varepsilon_3^2 = x_3^2 - 2 \frac{x_3^2}{n} + \frac{\sum x^2}{n^2},$$

$$\vdots \quad \vdots \quad \vdots \quad \vdots \quad \vdots \quad \vdots \quad \vdots \quad \vdots \quad \vdots \quad \vdots$$

$$\varepsilon_n^2 = x_n^2 - 2 \frac{x_n^2}{n} + \frac{\sum x^2}{n^2}.$$

После сложения этих уравнений, вводя для сокращения знак суммы $\sum$, находим:

$$\sum \varepsilon^2 = \sum x^2 - 2 \frac{\sum x^2}{n} + n \frac{\sum x^2}{n^2}.$$

Отсюда после приведения подобных членов находим:

$$\sum \varepsilon^2 = \left(1 - \frac{1}{n}\right) \sum x^2$$

или

$$\sum \varepsilon^2 = \frac{n-1}{n} \sum x^2. \quad (14)$$

На основании этого уравнения, а также на основании уравнения (7) мы можем для большего числа измерений одинаковой точности сделать следующие выводы:

1) Сумма квадратов отклонений ε результатов всех измерений от их среднего арифметического всегда меньше суммы квадратов точных значений ошибок x . Это следует из того, что коэффициент при $\sum x^2$ в правой части уравнений (14) всегда меньше единицы:

$$\frac{n-1}{n} < 1.$$

Но при увеличении числа измерений этот коэффициент возрастает и стремится к единице, поэтому мы вправе сказать, что при большом числе измерений сумма квадратов отклонений $\sum \varepsilon^2$ и сумма квадратов ошибок $\sum x^2$ практически оказываются равными.

2) Сумма квадратов точных значений ошибок $\sum x^2$ должна быть наименьшей, если сумма квадратов отклонений $\sum \varepsilon^2$ является наименьшей; справедливо, очевидно, и обратное заключение.

3) Точное значение X измеряемой величины отличается от среднего арифметического A результатов всех измерений на очень небольшую величину; она определяется из уравнения (7):

$$X - A = \frac{\sum x}{n}. \quad (15)$$

Числитель в правой части этого выражения остаётся малым при всяком числе измерений, а знаменатель возрастает при увеличении их числа. Отсюда следует, что при большом числе измерений среднее арифметическое их результатов чрезвычайно близко должно отвечать точному значению измеряемой величины.

3. Ошибки измерений—средняя квадратичная, вероятная и средняя. Кроме меры точности h , которая входит в основную формулу теории случайных ошибок, точность измерений принято определять ещё тремя показателями точности: средней квадратичной ошибкой измерений, вероятной ошибкой измерений и средней ошибкой измерений.

Эти показатели точности принято обозначать соответственно через σ , ρ и η . Для всех показателей точности

σ , ρ , η выведены выражения, которыми определяется их зависимость от величины случайных ошибок и числа измерений. Кроме того, все три показателя точности приведены в определённую зависимость от меры точности h . Далее необходимо иметь в виду, что можно определять показатели точности σ , ρ и η или для результатов отдельных измерений, или для их среднего арифметического. В обоих этих случаях выражения показателей точности оказываются различными. Наконец, следует ещё отметить, что в выражения показателей точности можно вводить или точные значения ошибок измерений x , или отклонения ε отдельных измерений от их среднего арифметического. В зависимости от этого выражения показателей точности в обоих этих случаях вновь оказываются несколько различными; например, в них может появиться или мера точности точных значений ошибок h , или мера точности отклонений H (стр. 166).

Основные выражения трёх показателей точности σ , ρ и η для отдельных измерений можно вывести из следующих соображений:

1) Средняя квадратичная ошибка σ отдельного измерения. Возьмём выражение, которым определяется вероятность совместного появления ряда ошибок $x_1, x_2, x_3, \dots, x_n$ при n измерениях одинаковой точности (стр. 160):

$$P = \left(\frac{h}{\sqrt{\pi}}\right)^n e^{-h^2(x_1^2 + x_2^2 + x_3^2 + \dots + x_n^2)} dx_1 dx_2 dx_3 \dots dx_n.$$

Определим то значение меры точности h для этого ряда измерений, при котором совместное появление ошибок $x_1, x_2, x_3, \dots, x_n$ оказывается наиболее вероятным. При этом условии вероятность P должна иметь наибольшее значение. Таким образом, мы должны считать h величиной переменной и найти для неё то значение, при котором вероятность P оказывается наибольшей. Для этого необходимо найти производную P по h и, положив её равной нулю, определить из полученного уравнения значение h . При этом значении h вероятность P достигает максимума.

Находим производную P по h :

$$\begin{aligned} \frac{dP}{dh} = & \left(\frac{h}{\sqrt{\pi}}\right)^n e^{-h^2(x_1^2 + x_2^2 + x_3^2 + \dots + x_n^2)} \times \\ & \times [-2h(x_1^2 + x_2^2 + x_3^2 + \dots + x_n^2)] + \\ & + e^{-h^2(x_1^2 + x_2^2 + x_3^2 + \dots + x_n^2)} n \left(\frac{h}{\sqrt{\pi}}\right)^{n-1} \cdot \frac{1}{\sqrt{\pi}}. \end{aligned}$$

Вынося за скобки общие множители, имеем:

$$\begin{aligned} \frac{dP}{dh} = & \frac{1}{\sqrt{\pi}} \cdot \left(\frac{h}{\sqrt{\pi}}\right)^{n-1} \cdot e^{-h^2(x_1^2 + x_2^2 + x_3^2 + \dots + x_n^2)} \times \\ & \times [-2h^2(x_1^2 + x_2^2 + x_3^2 + \dots + x_n^2) + n]. \end{aligned}$$

Положив эту производную равной нулю и сократив полученное уравнение на множитель, стоящий перед квадратными скобками, находим:

$$-2h^2(x_1^2 + x_2^2 + x_3^2 + \dots + x_n^2) + n = 0.$$

Отсюда получаем:

$$\frac{1}{2h^2} = \frac{x_1^2 + x_2^2 + x_3^2 + \dots + x_n^2}{n}$$

или окончательно:

$$\frac{1}{h\sqrt{2}} = \pm \sqrt{\frac{x_1^2 + x_2^2 + x_3^2 + \dots + x_n^2}{n}}.$$

Величину, стоящую в правой части последнего равенства, принято называть средней квадратичной ошибкой отдельного измерения σ . Таким образом, для определения средней квадратичной ошибки отдельного измерения имеем выражение:

$$\sigma = \pm \sqrt{\frac{x_1^2 + x_2^2 + x_3^2 + \dots + x_n^2}{n}} = \pm \sqrt{\frac{\sum x^2}{n}}. \quad (16)$$

Кроме того,

$$\sigma = \pm \frac{1}{h\sqrt{2}}. \quad (17)$$

В выражении (16) величины $x_1, x_2, x_3, \dots, x_n$ обозначают точные значения ошибок измерений, нам неизвестные. Для того чтобы найти выражение средней

квадратичной ошибки отдельного измерения через отклонение ϵ , воспользуемся тем же самым приёмом, при помощи которого мы вывели выражение (16). Вводя в вычисления величины $\epsilon_1, \epsilon_2, \epsilon_3, \dots, \epsilon_n$, мы вместо начального выражения (3) будем иметь такое выражение:

$$P' = \left(\frac{H}{\sqrt{\pi}} \right)^n e^{-H^2 (\epsilon_1^2 + \epsilon_2^2 + \epsilon_3^2 + \dots + \epsilon_n^2)},$$

где H определяется формулой (10).

Выполняя над этим выражением все преобразования, которые мы проделали с выражением (3), получаем:

$$\frac{1}{H\sqrt{2}} = \pm \sqrt{\frac{\epsilon_1^2 + \epsilon_2^2 + \epsilon_3^2 + \dots + \epsilon_n^2}{n}}.$$

Отсюда, принимая во внимание выражение (10), находим:

$$\frac{1}{h\sqrt{2}} \cdot \frac{1}{\sqrt{\frac{n}{n-1}}} = \pm \sqrt{\frac{\epsilon_1^2 + \epsilon_2^2 + \epsilon_3^2 + \dots + \epsilon_n^2}{n}}$$

или

$$\frac{1}{h\sqrt{2}} = \pm \sqrt{\frac{\epsilon_1^2 + \epsilon_2^2 + \epsilon_3^2 + \dots + \epsilon_n^2}{n-1}}.$$

Из этого выражения на основании уравнения (17) получаем окончательно:

$$\sigma = \pm \sqrt{\frac{\epsilon_1^2 + \epsilon_2^2 + \epsilon_3^2 + \dots + \epsilon_n^2}{n-1}} = \pm \sqrt{\frac{\sum \epsilon^2}{n-1}}. \quad (16')$$

Таким образом, мы видим, что если среднюю квадратичную ошибку выражать через отклонения ϵ результатов отдельных измерений от их среднего арифметического, то в знаменатель под знаком радикала входит число измерений n , уменьшенное на единицу. Этот вывод находится в полном соответствии с формулой (14), которая даёт возможность из уравнения (16) непосредственно получить выражение (16') прямой заменой суммы квадратов ошибок $\sum x^2$ суммой квадратов отклонений $\sum \epsilon^2$, умноженной на обратную величину коэффициента при $\sum x^2$ в формуле (14).

2) Вероятная ошибка ρ отдельного измерения. Вероятной ошибкой ρ отдельного измерения называют такую ошибку, которая делит все случайные ошибки данного ряда n измерений на две равные по числу ошибок части, так что в одной части находятся $\frac{n}{2}$ случайных ошибок, больших вероятной ошибки ρ , а в другой части $\frac{n}{2}$ случайных ошибок, меньших её. Иначе можно сказать, что вероятность того, что некоторая случайная ошибка лежит в пределах от $-\rho$ до $+\rho$, должна равняться половине, так как число случайных ошибок, которые лежат в этих пределах, равно половине общего числа всех случайных ошибок.

Таким образом, на основании выражения (27), стр. 151 мы можем написать:

$$\frac{h}{\sqrt{\pi}} \int_{-\rho}^{+\rho} e^{-h^2 x^2} dx = \frac{1}{2}$$

или

$$\frac{h}{\sqrt{\pi}} \int_0^{\rho} e^{-h^2 x^2} dx = \frac{1}{4}.$$

Отсюда получаем:

$$\int_0^{\rho} e^{-h^2 x^2} h dx = \frac{\sqrt{\pi}}{4}.$$

Для того чтобы из этого выражения определить величину ρ , сделаем подстановку, полагая:

$$hx = z \quad \text{и} \quad h dx = dz.$$

Вводя это новое переменное под знак интеграла и изменяя соответствующим образом пределы интегрирования, находим:

$$\int_0^{h\rho} e^{-z^2} dz = \frac{\sqrt{\pi}}{4}. \quad (18)$$

Введём обозначение $h\rho = t$. Имеем:

$$\int_0^t e^{-z^2} dz = \frac{\sqrt{\pi}}{4}.$$

Этот интеграл нам уже приходилось вычислять при решении уравнения (30), стр. 153. Мы пользовались при этом методом приближённого вычисления определённых интегралов, применяя разложение функции e^x в бесконечный ряд.

Вновь пользуясь тем же приёмом, мы можем написать:

$$\int_0^t e^{-z^2} dz = \int_0^t \left(1 - z^2 + \frac{z^4}{2!} - \frac{z^6}{3!} + \frac{z^8}{4!} - \frac{z^{10}}{5!} + \dots \right) dz = \frac{\sqrt{\pi}}{4}.$$

Выполняя интегрирование и ограничиваясь при этом только написанными членами ряда, получаем такое уравнение:

$$t - \frac{t^3}{3} + \frac{t^5}{10} - \frac{t^7}{42} + \frac{t^9}{216} - \frac{t^{11}}{1320} - \frac{\sqrt{\pi}}{4} = 0.$$

Это уравнение можно решить только приближённо. Его действительный корень, вычисленный с точностью до четвёртого десятичного знака, оказывается равным:

$$t = 0,4769. \quad (19)$$

Подставляя это значение в формулу (18), находим:

$$\rho = \pm 0,4769 \cdot \frac{1}{h}.$$

Из этой формулы на основании выражения (17) определяем зависимость между ρ и σ :

$$\rho = 0,4769 \cdot \sqrt{2} \cdot \sigma.$$

Выполняя вычисления в правой части этого равенства, получаем с точностью до четвёртого десятичного знака:

$$\rho = 0,6745\sigma. \quad (20)$$

Отсюда на основании формулы (16) имеем:

$$\begin{aligned} \rho &= \pm 0,6745 \sqrt{\frac{x_1^2 + x_2^2 + x_3^2 + \dots + x_n^2}{n}} = \\ &= \pm 0,6745 \sqrt{\frac{\sum x^2}{n}}. \end{aligned} \quad (21)$$

Кроме того, на основании выражения (16') находим также:

$$\rho = \pm 0,6745 \sqrt{\frac{\varepsilon_1^2 + \varepsilon_2^2 + \varepsilon_3^2 + \dots + \varepsilon_n^2}{n-1}} = \pm 0,6745 \sqrt{\frac{\sum \varepsilon^2}{n-1}}. \quad (21')$$

Этими выражениями устанавливается зависимость между средней квадратичной и вероятной ошибками отдельного измерения; коэффициентом перехода от σ к ρ служит величина $0,4769\sqrt{2}$, равная с точностью до четвёртого десятичного знака 0,6745, которую очень часто принимают приближённо равной $\frac{2}{3}$, так что последнюю формулу можно написать ещё в таком виде:

$$\rho = \pm \frac{2}{3} \sqrt{\frac{\varepsilon_1^2 + \varepsilon_2^2 + \varepsilon_3^2 + \dots + \varepsilon_n^2}{n-1}} = \pm \frac{2}{3} \sqrt{\frac{\sum \varepsilon^2}{n-1}}.$$

3) Средняя ошибка η отдельного измерения. Средней ошибкой η отдельного измерения называют среднее арифметическое абсолютных величин всех случайных ошибок данного ряда n измерений.

Таким образом, мы можем написать:

$$\eta = \pm \frac{|x_1| + |x_2| + |x_3| + \dots + |x_n|}{n} = \pm \frac{\sum |x|}{n}. \quad (22)$$

Эта формула не устанавливает зависимости средней ошибки отдельного измерения от меры точности h и от остальных двух показателей точности σ и ρ . Для того чтобы определить эту зависимость, воспользуемся простым, хотя и не вполне точным выводом. Допустим по-прежнему, что сделано n измерений некоторой величины. Все измерения имеют одинаковую точность; точные значения их ошибок будем обозначать по-прежнему через x с соответствующими индексами.

Вероятность некоторой случайной ошибки, равной x , согласно формуле (26), стр. 151, определяется выражением

$$dP = \frac{h}{\sqrt{\pi}} e^{-h^2 x^2} dx.$$

Вероятное число случайных ошибок той же величины во всём ряде n измерений должно быть в n раз больше. Поэтому, обозначая это число через N , можно написать:

$$N = n \frac{h}{\sqrt{\pi}} e^{-h^2 x^2} dx.$$

Сумма абсолютных величин всех таких ошибок должна определяться их числом N , умноженным на абсолютную величину ошибки, т. е. на $|x|$. Таким образом, обозначая эту сумму через S , мы можем написать:

$$S = |x| \cdot n \frac{h}{\sqrt{\pi}} e^{-h^2 x^2} dx.$$

Наконец, для того чтобы определить сумму абсолютных величин всех ошибок ряда n измерений, т. е. для того, чтобы определить числитель в правой части формулы (22), надо последнее выражение интегрировать в пределах от $-\infty$ до $+\infty$. Таким образом, можно написать:

$$\sum |x| = \int_{-\infty}^{+\infty} |x| \cdot n \frac{h}{\sqrt{\pi}} e^{-h^2 x^2} dx$$

или

$$\sum |x| = \frac{2n}{h\sqrt{\pi}} \int_0^{\infty} e^{-h^2 x^2} h^2 x dx.$$

Этот интеграл можно непосредственно вычислить; действительно, напомним последнее выражение в таком виде:

$$\sum |x| = -\frac{n}{h\sqrt{\pi}} \int_0^{\infty} e^{-(hx)^2} \cdot (-2h^2 x dx).$$

Мы видим, что под знаком интеграла стоит дифференциал выражения $e^{-h^2 x^2}$. Вследствие этого, выполняя

интегрирование, находим:

$$\sum |x| = -\frac{n}{h\sqrt{\pi}} \left|_0^{\infty} e^{-h^2 x^2} \right|$$

или

$$\sum |x| = \frac{n}{h\sqrt{\pi}}.$$

Вставляя это значение $\sum |x|$ в выражение (22), находим:

$$\eta = \pm \frac{1}{h\sqrt{\pi}} \quad (23)$$

и на основании выражения (17) определяем зависимость между η и σ :

$$\eta = \sqrt{\frac{2}{\pi}} \cdot \sigma. \quad (24)$$

Выполняя вычисления в правой части этого равенства, находим с точностью до четвёртого десятичного знака:

$$\eta = 0,7979 \sigma. \quad (25)$$

Отсюда на основании выражений (16) и (16') получаем:

$$\begin{aligned} \eta &= \pm 0,7979 \sqrt{\frac{x_1^2 + x_2^2 + x_3^2 + \dots + x_n^2}{n}} = \\ &= \pm 0,7979 \sqrt{\frac{\sum x^2}{n}}, \end{aligned} \quad (26)$$

$$\begin{aligned} \eta &= \pm 0,7979 \sqrt{\frac{\epsilon_1^2 + \epsilon_2^2 + \epsilon_3^2 + \dots + \epsilon_n^2}{n-1}} = \\ &= \pm 0,7979 \sqrt{\frac{\sum \epsilon^2}{n-1}}. \end{aligned} \quad (26')$$

Этими выражениями устанавливается зависимость между средней квадратичной и средней ошибкой отдельного измерения. Коэффициентом для перехода от σ к η служит величина $\sqrt{\frac{2}{\pi}}$, равная с точностью до четвёртого

десятичного знака 0,7979, которую очень часто принимают приближённо равной 0,8, так что последнюю формулу можно написать ещё в таком виде:

$$\eta = \pm 0,8 \sqrt{\frac{\epsilon_1^2 + \epsilon_2^2 + \epsilon_3^2 + \dots + \epsilon_n^2}{n-1}} = \pm 0,8 \sqrt{\frac{\sum \epsilon^2}{n-1}}. \quad (27)$$

Зависимость между ошибками σ и η и мерой точности h очень часто выражают формулами с числовыми коэффициентами. Вычисляя эти коэффициенты согласно формулам (17) и (23) и присоединяя сюда формулу, определяющую зависимость ошибки ρ от меры точности h , получаем:

$$\begin{aligned} \sigma &= \pm \frac{1}{h \sqrt{2}} = \pm 0,7071 \cdot \frac{1}{h}, \\ \rho &= \pm 0,4769 \cdot \frac{1}{h}, \\ \eta &= \pm \frac{1}{h \sqrt{\pi}} = \pm 0,5642 \cdot \frac{1}{h}. \end{aligned} \quad (28)$$

Что же касается взаимной зависимости трёх показателей точности σ , ρ и η , то коэффициенты для перехода от одного показателя к другому можно вычислить из формул (28). Величина этих коэффициентов, вычисленная с точностью также до четвёртого десятичного знака, приведена в таблице IV.

Таблица IV

	σ	ρ	η
σ	1,0000	1,4826	1,2533
ρ	0,6745	1,0000	0,8453
η	0,7979	1,1829	1,0000

При вычислениях небольшой точности обычно ограничиваются приближёнными значениями этих коэффициентов, отбрасывая два, а иногда и три последних знака.

4. Геометрическое значение средней квадратичной, вероятной и средней ошибок отдельных измерений. Выше мы разобрали геометрическое значение меры точности h в основной формуле случайных ошибок. Также нетрудно выяснить и геометрический смысл всех трёх показателей точности σ , ρ и η , т. е. средней квадратичной, вероятной и средней ошибок отдельных измерений. Прежде всего следует обратить внимание на то, что все три показателя точности представляют собой некоторые ошибки измерений, и так как величину ошибок мы откла-

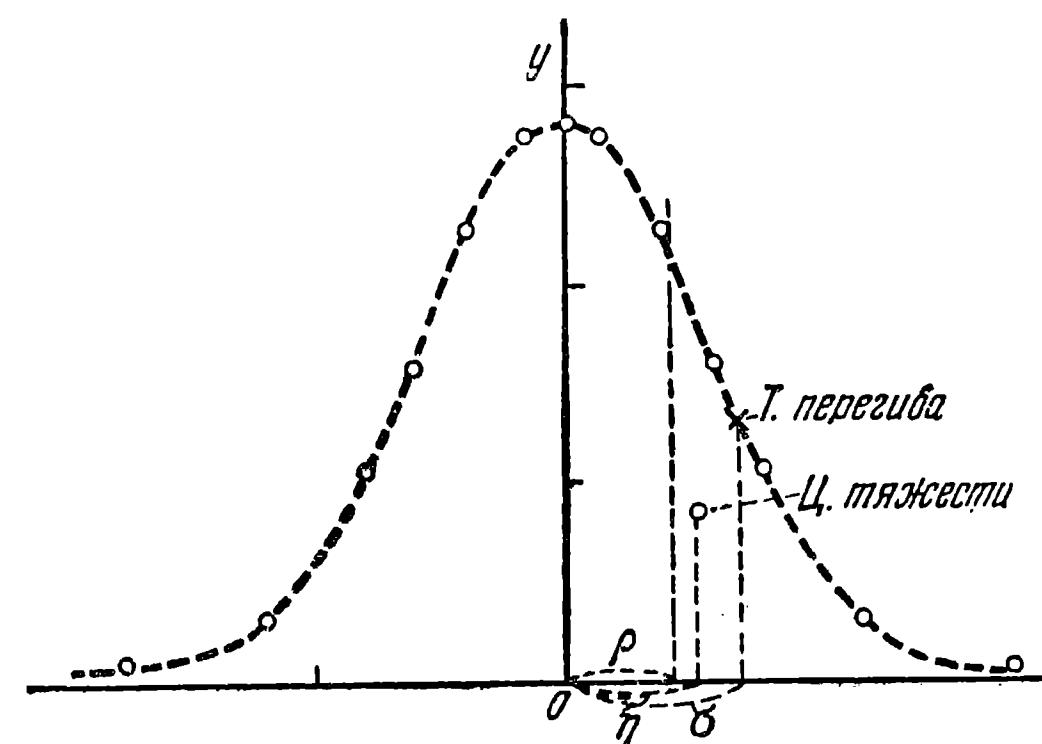


Рис. 10.

дываем по оси абсцисс, то все три показателя точности должны геометрически изображаться некоторыми отрезками оси абсцисс.

Попробуем определить эти отрезки. Для этого возьмём кривую ошибок, мера точности h которой может быть какой угодно (рис. 10). Уравнение кривой имеет вид

$$y = \frac{h}{\sqrt{\pi}} e^{-h^2 x^2}.$$

Такая кривая должна иметь две точки перегиба, которые лежат симметрично по отношению к оси ординат. Это следует из того, что в своей верхней части кривая обращена выпуклой стороной кверху, а при приближении

к оси абсцисс выпуклая сторона обеих ветвей кривой оказывается обращённой книзу.

Определим абсциссы точек перегиба кривой ошибок. Для этого находим вторую производную y по x :

$$\frac{dy}{dx} = -\frac{2h^3}{\sqrt{\pi}} x \cdot e^{-h^2 x^2},$$

$$\frac{d^2 y}{dx^2} = \frac{4h^5 x^2}{\sqrt{\pi}} e^{-h^2 x^2} - \frac{2h^3}{\sqrt{\pi}} e^{-h^2 x^2}$$

или

$$\frac{d^2 y}{dx^2} = -\frac{2h^3}{\sqrt{\pi}} e^{-h^2 x^2} (1 - 2h^2 x^2).$$

Полагая вторую производную равной нулю, находим:

$$1 - 2h^2 x^2 = 0.$$

Из этого уравнения определяем x :

$$x = \pm \frac{1}{h\sqrt{2}}.$$

Эти значения x определяют абсциссы точек перегиба на левой и на правой ветвях кривой ошибок. Сравнивая это выражение с формулой (17), мы видим, что той же величиной определяется значение средней квадратичной ошибки σ отдельного измерения. Таким образом, мы вправе сделать следующий вывод:

Средняя квадратичная ошибка σ отдельного измерения геометрически определяется абсциссами точек перегиба на обеих половинах кривой ошибок.

По отношению к вероятной ошибке отдельного измерения мы можем непосредственно утверждать, что:

Вероятная ошибка ρ отдельного измерения геометрически определяется абсциссами тех точек, ординаты которых делят на две равные части площадь, ограниченную кривой ошибок, осью ординат и осью абсцисс в её положительной или отрицательной части.

Этот вывод непосредственно следует из самого определения вероятной ошибки отдельного измерения.

Наконец, для того, чтобы определить геометрический смысл средней ошибки отдельного измерения, найдём абсциссы центра тяжести каждой половины площади, ограниченной кривой ошибок и осью абсцисс.

Если дано уравнение кривой в декартовых координатах:

$$y = f(x),$$

то абсцисса x_0 центра тяжести площади, ограниченной этой кривой и осями координат, определяется выражением

$$x_0 = \frac{\int xy \, dx}{\int y \, dx}.$$

Применяя это выражение к данному случаю, находим:

$$x_0 = \frac{\int_0^\infty x \frac{h}{\sqrt{\pi}} e^{-h^2 x^2} dx}{\int_0^\infty \frac{h}{\sqrt{\pi}} e^{-h^2 x^2} dx}$$

или

$$x_0 = \frac{\frac{1}{\sqrt{\pi}} \int_0^\infty h e^{-h^2 x^2} x \, dx}{\frac{h}{\sqrt{\pi}} \int_0^\infty e^{-h^2 x^2} dx}.$$

Интеграл, стоящий в знаменателе этого выражения, равен $\frac{1}{2}$, так как он равен половине интеграла вероятностей при бесконечном возрастании его пределов (стр. 148). Что же касается интеграла в числителе этого выражения, то мы можем его привести к виду

$$\frac{1}{\sqrt{\pi}} \int_0^\infty h e^{-h^2 x^2} x \, dx = -\frac{1}{2h\sqrt{\pi}} \int_0^\infty e^{-h^2 x^2} \cdot (-2h^2 x \, dx).$$

Мы получили интеграл, который уже вычисляли, определяя зависимость между средней ошибкой и мерой

точности (стр. 178). Таким образом, выражение, определяющее абсциссу x_0 центра тяжести, после интегрирования получает такой вид:

$$x_0 = \frac{\frac{1}{2h\sqrt{\pi}}}{\frac{1}{2}} = \frac{1}{h\sqrt{\pi}}.$$

Сравнивая это выражение с выражением (23), мы можем сделать такой вывод:

Средняя ошибка η отдельного измерения геометрически определяется абсциссами центров тяжести двух площадей, ограниченных обеими ветвями кривой ошибок и осями координат.

Относительные размеры и положение всех трёх показателей точности даны на рис. 10. Такое же построение можно сделать и для отрицательной части оси абсцисс, как это следует на основании двойного знака у показателей точности σ , ρ и η .

Следует обратить внимание на то, что из трёх ошибок σ , ρ и η наибольшую величину имеет средняя квадратичная ошибка σ , а наименьшую — вероятная ошибка ρ .

5. Ошибки среднего арифметического. Характеризуя окончательную точность измерений, принято указывать не ошибку отдельных измерений, а ошибку их среднего арифметического, т. е. окончательного результата измерений. Для того чтобы вывести выражения, которые определяют ошибки среднего арифметического по ошибкам отдельных измерений, необходимо предварительно рассмотреть теорему об определении средней квадратичной ошибки линейной функции ряда независимых величин, средние квадратичные ошибки которых нам известны. Эту теорему, которая находит и самостоятельное применение, мы рассмотрим сначала для двух независимых величин, а затем перейдём к общему случаю.

Допустим, что мы произвели измерения двух величин B_1 и B_2 , которые определяются независимо одна от другой

на основании отдельных измерений. Для величины B_1 мы произвели n_1 измерений, а для величины B_2 произвели n_2 измерений. Все измерения каждого ряда имеют одинаковую точность. Обозначим точные значения ошибок измерений величин B_1 и B_2 соответственно через $x'_1, x'_2, x'_3, \dots, x'_{n_1}$ и через $x''_1, x''_2, x''_3, \dots, x''_{n_2}$, а точные значения средних квадратичных ошибок отдельного измерения величин B_1 и B_2 соответственно через σ_1 и σ_2 . Величина последних на основании формулы (16) определяется выражениями

$$\sigma_1 = \pm \sqrt{\frac{\sum x'^2}{n_1}}, \quad \sigma_2 = \pm \sqrt{\frac{\sum x''^2}{n_2}}. \quad (29)$$

Определим среднюю квадратичную ошибку суммы величин B_1 и B_2 в зависимости от средних квадратичных ошибок σ_1 и σ_2 отдельных измерений этих величин. Введём такие обозначения: сумму величин B_1 и B_2 обозначим через B ; число ошибок величины B , которое должно зависеть от числа ошибок n_1 и n_2 величин B_1 и B_2 , обозначим через n ; наконец, обозначим величины ошибок B через $x_1, x_2, x_3, \dots, x_n$ и среднюю квадратичную ошибку величины B через σ .

Мы допускаем, что средняя квадратичная ошибка суммы величин B_1 и B_2 определяется выражением, аналогичным формулам (29), т. е. полагаем, что

$$\sigma = \pm \sqrt{\frac{\sum x^2}{n}}, \quad (30)$$

и для определения σ найдём величины $\sum x^2$ и n .

Ошибки величины B на основании теоремы об абсолютной ошибке суммы (стр. 59) должны равняться суммам ошибок величин B_1 и B_2 , т. е. суммам величин $x'_1, x'_2, x'_3, \dots, x'_{n_1}$ и $x''_1, x''_2, x''_3, \dots, x''_{n_2}$, взятых попарно, причём, так как все ошибки того и другого ряда измерений являются вполне независимыми, то каждая из ошибок первого ряда может появиться совместно с каждой из ошибок второго ряда. Отсюда следует, что все возможные ошибки величины B должны определяться

получает такой вид:

$$\sigma = \pm \sqrt{b_1^2 \sigma_1^2 + b_2^2 \sigma_2^2}. \quad (33')$$

Наконец, в общем случае величина B может быть линейной функцией ряда независимых величин $B_1, B_2, B_3, \dots, B_m$, т. е. может определяться выражением

$$B = b_1 B_1 + b_2 B_2 + b_3 B_3 + \dots + b_m B_m,$$

где $b_1, b_2, b_3, \dots, b_m$ обозначают некоторые постоянные коэффициенты. Теорема оказывается справедливой и для этого случая, т. е. в формуле под знаком корня входит также сумма квадратов средних квадратичных ошибок отдельных измерений величин $B_1, B_2, B_3, \dots, B_m$, т. е. величин $\sigma_1^2, \sigma_2^2, \sigma_3^2, \dots, \sigma_m^2$, и каждая из этих величин умножается на квадрат соответствующего коэффициента. Таким образом, в этом общем случае формула (33) получает такой вид:

$$\sigma = \pm \sqrt{b_1^2 \sigma_1^2 + b_2^2 \sigma_2^2 + b_3^2 \sigma_3^2 + \dots + b_m^2 \sigma_m^2}. \quad (33'')$$

Воспользуемся последней формулой для того, чтобы вывести выражение средней квадратичной ошибки среднего арифметического в зависимости от средней квадратичной ошибки отдельных измерений. Для этого берём ряд n измерений, результаты которых обозначаем через $a_1, a_2, a_3, \dots, a_n$ и их среднее арифметическое представим в таком виде:

$$A = \frac{1}{n} a_1 + \frac{1}{n} a_2 + \frac{1}{n} a_3 + \dots + \frac{1}{n} a_n.$$

В этом выражении величину A можно рассматривать как линейную функцию величин $a_1, a_2, a_3, \dots, a_n$. Последние величины являются независимыми и имеют одну и ту же среднюю квадратичную ошибку σ , которая определяется формулами (16) и (16'). Поэтому, обозначая среднюю квадратичную ошибку среднего арифметического через σ_A и применяя формулу (33''), мы можем написать:

$$\sigma_A = \pm \sqrt{n \cdot \left(\frac{1}{n}\right)^2 \sigma^2} = \pm \frac{\sigma}{\sqrt{n}}. \quad (35)$$

Подставляя в это выражение значения σ из формул (16) и (16'), получаем:

$$\sigma_A = \pm \sqrt{\frac{x_1^2 + x_2^2 + x_3^2 + \dots + x_n^2}{n}} = \pm \frac{\sqrt{\sum x^2}}{n} \quad (36)$$

и

$$\sigma_A = \pm \sqrt{\frac{\epsilon_1^2 + \epsilon_2^2 + \epsilon_3^2 + \dots + \epsilon_n^2}{n(n-1)}} = \pm \sqrt{\frac{\sum \epsilon^2}{n(n-1)}}. \quad (36')$$

Вероятную ошибку среднего арифметического ρ_A мы получим, умножая его среднюю квадратичную ошибку σ_A на выведенный выше коэффициент 0,6745. Таким образом, мы можем написать:

$$\rho_A = \pm 0,6745 \sqrt{\frac{\sum \epsilon^2}{n(n-1)}}. \quad (37)$$

При этом уравнение (35) остаётся в силе и для вероятных ошибок, т. е. вероятная ошибка среднего арифметического находится в том отношении к вероятной ошибке отдельного измерения, которое определяется формулой (35). Таким образом, можно написать:

$$\rho_A = \frac{\rho}{\sqrt{n}}. \quad (35')$$

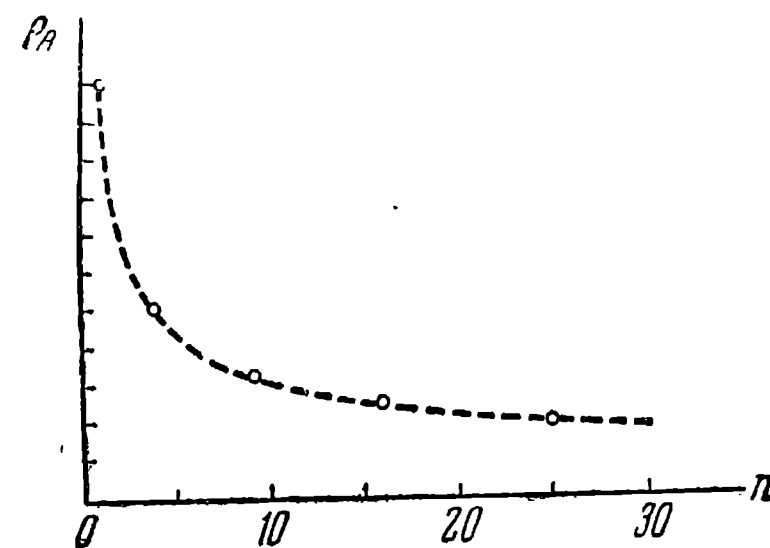


Рис. 11.

Если уравнение (35) или (35') изобразить графически, откладывая по оси абсцисс число измерений n , а по оси ординат ошибку среднего арифметического, например ρ_A (рис. 11), то мы видим, что при увеличении числа измерений вероятная ошибка среднего арифметического (так же как и средняя квадратичная) непрерывно уменьшается, но это уменьшение, начиная с некоторого числа измерений, равного приблизительно десяти или пятнадцати, становится очень незначительным. Отсюда следует, что в целях уменьшения ошибок среднего арифметического не рационально идти по пути чрезмерного увеличения

числа измерений; последнее может быть оправдано только общими требованиями достаточно обоснованного применения теории случайных ошибок. В целях же уменьшения ошибок окончательного результата измерений, если это оказывается необходимым, следует увеличивать меру их точности h , т. е. применять более точные методы.

Вопрос о том, каким из показателей точности следует пользоваться при оценке точности окончательного результата измерений, нельзя считать вполне решённым, и в различных научных школах применяют различные показатели точности. Легче всего вычисляется средняя ошибка отдельного измерения, однако при тщательной обработке результатов измерений, если они выполнены с большой точностью, обыкновенно определяют квадратичную или вероятную ошибку среднего арифметического, хотя это сопровождается более сложными вычислениями.

6. Взвешенные измерения и исправление результатов.

При выводе всех основных формул, приведенных выше, мы всегда предполагаем, что в данном ряде n измерений все отдельные измерения $a_1, a_2, a_3, \dots, a_n$ имеют одинаковую меру точности h , т. е. что нет никаких оснований отдавать предпочтение одним каким-либо измерениям этого ряда по сравнению с остальными. Но иногда при выполнении измерений может казаться, что некоторые из них протекают особенно удачно, например, при более благоприятных внешних условиях, — температурные условия могли оставаться особенно постоянными или совершенно не замечалось влияния случайных индукционных токов в общей электрической сети и т. п. В таких случаях экспериментатор часто рассматривает результаты этих измерений как наиболее ценные или наиболее надёжные; вследствие этого у него может возникнуть вполне понятное желание выделить эти, повидимому, наиболее надёжные измерения и при вычислении среднего арифметического усилить их влияние. Из таких соображений возникло предположение о возможности исправлять окончательные результаты измерений при помощи такого приёма: при вычислении среднего арифметического вводили результаты отдельных наиболее надёжных измерений с некоторы-

ми множителями, большими единицы, которые выбирались тем большими, чем надёжнее казалось данное измерение. Нетрудно видеть, что такой метод исправления окончательного результата измерений является совершенно неприемлемым. Выбирая коэффициенты для результатов отдельных измерений или, как было принято говорить, приписывая отдельным измерениям определённый вес, мы могли бы установить его величину на основании только чисто субъективной оценки, так как подтвердить каким-либо объективным расчётом надёжность или ценность отдельных измерений, вообще говоря, не представляется возможным. Вследствие этого теория взвешенных измерений в такой форме не может быть предметом серьёзного обсуждения; такая теория могла бы не исправить результаты измерений, а, наоборот, сделать их менее вероятными, и это является настолько несомненным, что от подобного приёма в настоящее время почти совершенно отказались.

Вместе с тем теория взвешенных измерений, но в совершенно иной, математически разработанной форме, находит применение, например, при определении наиболее вероятного значения какой-либо величины, для которой было получено несколько различных значений из нескольких независимых рядов измерений, имеющих различную меру точности. Подробное изложение этих вопросов, как относительно редких в практике обычных лабораторных измерений, мы приводить не будем и ограничимся только несколькими элементарными замечаниями, которые могут быть полезными в дальнейшем.

В основе математической теории взвешенных измерений лежит точное определение веса измерений. Это можно обосновать следующими соображениями.

Если результаты n измерений мы обозначим попрежнему через $a_1, a_2, a_3, \dots, a_n$, но допустим, что все они имеют различные меры точности $h_1, h_2, h_3, \dots, h_n$, то последние величины должны появиться в правых частях уравнений (2). Условие минимума мы должны в этом случае наложить на величину:

$$x_1^2 h_1^2 + x_2^2 h_2^2 + x_3^2 h_3^2 + \dots + x_n^2 h_n^2 = \sum x^2 h^2.$$

Отсюда видно, что в случае неравноточных измерений принцип наименьших квадратов требует, чтобы при вычислении суммы квадратов ошибок x каждая из них была предварительно помножена на меру точности соответствующего измерения. Этими множителями при x^2 в последнем выражении и определяются веса измерений. Таким образом, обозначая веса измерений соответственно через $w_1, w_2, w_3, \dots, w_n$, мы можем условие минимума написать в таком виде:

$$x_1^2 w_1 + x_2^2 w_2 + x_3^2 w_3 + \dots + x_n^2 w_n = \sum x^2 w = \text{мин.}$$

Из этого условия находим прежним приёмом (стр. 161), что наиболее вероятное или наилучшее значение A_w для ряда n неравноточных измерений определяется таким выражением:

$$A_w = \frac{a_1 w_1 + a_2 w_2 + a_3 w_3 + \dots + a_n w_n}{w_1 + w_2 + w_3 + \dots + w_n} = \frac{\sum a w}{\sum w}. \quad (38)$$

Величина A_w , определяемая формулой (38), получила название среднего взвешенного n измерений с различными мерами точности.

Следует указать на то, что условие минимума остаётся в силе (стр. 161), если от ошибок измерений x мы переходим к отклонениям ϵ отдельных измерений от среднего взвешенного. Вследствие этого во всех формулах, куда входит сумма квадратов отклонений, при каждой из величин ϵ^2 появляется множитель, равный весу данного измерения. Так, для квадратичной ошибки w^{σ_A} среднего взвешенного мы находим такое выражение:

$$w^{\sigma_A} = \pm \sqrt{\frac{\epsilon_1^2 w_1 + \epsilon_2^2 w_2 + \epsilon_3^2 w_3 + \dots + \epsilon_n^2 w_n}{(n-1)(w_1 + w_2 + w_3 + \dots + w_n)}} = \\ = \pm \sqrt{\frac{\sum \epsilon^2 w}{(n-1) \sum w}}. \quad (39)$$

Точно так же для средней квадратичной ошибки w^{σ_k} отдельного измерения, вес которого равен w_k , получается

такое выражение:

$$w^{\sigma_k} = \pm \sqrt{\frac{\epsilon_1^2 w_1 + \epsilon_2^2 w_2 + \epsilon_3^2 w_3 + \dots + \epsilon_n^2 w_n}{(n-1) w_k}} = \\ = \pm \sqrt{\frac{\sum \epsilon^2 w}{(n-1) w_k}}. \quad (39')$$

Необходимо также обратить внимание на следующее. Принимая, что веса измерений пропорциональны вторым степеням их мер точности, мы согласно формулам (28) должны одновременно допустить, что веса измерений обратно пропорциональны квадратам показателей точности σ, ρ и η . Так, для средней квадратичной ошибки отдельного измерения σ можно написать:

$$\frac{w_1}{w_2} = \frac{\sigma_2^2}{\sigma_1^2},$$

где w_1 и w_2 обозначают веса двух измерений, средние квадратичные ошибки которых равны соответственно σ_1 и σ_2 .

Теория взвешенных измерений при обычных лабораторных работах, как уже было сказано, не находит применения, так как мы, вообще говоря, всегда предполагаем, что все измерения данной серии имеют одинаковую меру точности. Приступая к обработке результатов измерений, нам приходится прежде всего определить их среднее арифметическое. При его вычислении мы должны пользоваться только общей формулой (5), не вводя в неё никаких произвольных коэффициентов.

Затем необходимо найти величины ϵ и отсюда вычислить среднюю ошибку отдельного измерения η . Эти вычисления позволяют производить общую оценку результатов отдельных измерений. Действительно, при правильно выполненных измерениях значения ϵ не должны быть много больше значения η . Это следует из того, что ошибки отдельных измерений при тщательном выполнении их не должны выходить далеко за те пределы, которые отвечают общей точности данных измерений (стр. 25). Но иногда, вычислив ϵ и η , мы обнаруживаем, что некоторые из ϵ имеют ненормально большую величину: в шесть, семь и более раз превосходят значение η . Такой ряд

измерений мы должны признать неудачным. Никаких исправлений в эти измерения мы вносить не можем, мы должны отказаться от этих результатов вообще и весь ряд измерений повторить ещё раз более тщательно.

Только в одном случае допускают исключение и не повторяют измерений вторично: если число измерений n достаточно велико и ненормально большую величину обнаруживают одно или в крайнем случае два значения ϵ и если, кроме того, повторение измерений встречает какие-либо серьёзные затруднения. Если только одно или два ϵ ненормально велики, то есть основания считать, что данный ряд измерений мы выполнили удачно, но одно или два измерения оказались ошибочными, — эти измерения следует отбросить. После этого необходимо найти новые значения: среднего арифметического, величин ϵ и средней ошибки отдельного измерения η . Такое исправление ряда уже законченных измерений является единственным, которое допускается, но и то в виде редких исключений.

Переходя к вычислению окончательного результата измерений, необходимо выбрать определённый показатель точности. Если мы останавливаемся на квадратичной ошибке σ_A или на вероятной ошибке ρ_A , как это обычно и имеет место, то значение окончательного результата измерений A определяется формулами

$$A = \frac{\sum a}{n} \pm \sqrt{\frac{\sum \epsilon^2}{n(n-1)}}, \quad (40)$$

$$A = \frac{\sum a}{n} \pm \frac{2}{3} \sqrt{\frac{\sum \epsilon^2}{n(n-1)}}. \quad (40')$$

Первый член в правой части этих формул определяет наиболее вероятное значение измеряемой величины, а второй член даёт величину соответственно его квадратичной и его вероятной ошибок. Очевидно, что определение окончательного значения измеряемой величины требует дополнительных вычислений, так как хотя значение среднего арифметического нам уже известно из предшествующих вычислений, но для определения σ_A или ρ_A требуется новая вычислительная работа.

Для того чтобы познакомиться на практике с выполнением подобных вычислений, рассмотрим несколько примеров.

Пример 1. Сопротивление R катушки измерялось при помощи мостика. Измерения были повторены десять раз; их результаты, выраженные в омах, оказались таковы:

6,270 6,271 6,276 6,278 6,277

6,277 6,273 6,272 6,275 6,274

Найдём из этих данных наиболее вероятное значение R и его квадратичную ошибку.

Все необходимые для этого вычисления приведены в таблице V, в которой указан и их окончательный результат, т. е. значение R и его квадратичная ошибка.

Таблица V

R			ε		ε²	
1	6,270	0,017	+0,004	$\sum \epsilon^+ = 0,010$ $\sum \epsilon^- = -0,013$ $\sum \epsilon = -0,003$	16·10 ⁻⁶	39·10 ⁻⁶
2	6,277		-0,003		9·10 ⁻⁶	
3	6,271		+0,003		9·10 ⁻⁶	
4	6,273		+0,001		1·10 ⁻⁶	
5	6,276		-0,002		4·10 ⁻⁶	
6	6,272	0,026	+0,002		4·10 ⁻⁶	30·10 ⁻⁶
7	6,278		-0,004		16·10 ⁻⁶	
8	6,275		-0,001		1·10 ⁻⁶	
9	6,277		-0,003		9·10 ⁻⁶	
10	6,274		0		0	
		0,043	$\sum \epsilon = 0,023 \quad \eta = \pm 0,002$		$\sum \epsilon^2 = 0,000069$	
$\frac{\sum R}{10} = 6,274 (3)$			$\sigma = \pm \sqrt{\frac{0,000069}{9}} = \pm 0,0028,$ $\sigma_R = \pm \frac{0,0028}{\sqrt{10}} = \pm 0,0009$			
$R = (6,274 \pm 0,001) \text{ ом}$						

По отношению к этой таблице необходимо сделать несколько замечаний:

1) Во всех значениях R , расположенных в первом столбце таблицы, три первые цифры остаются без изменения, вследствие этого при нахождении $\sum R$ достаточно сложить только последние цифры R . При большом числе слагаемых следует выписывать отдельные итоги, например, через каждые пять слагаемых, как это показано во втором столбце таблицы (0,017 и 0,026), и затем находить их общую сумму (0,043). Отсюда определяем последнюю цифру $\sum R$ и выписываем значение среднего арифметического, указывая в нём одну лишнюю цифру (стр. 77), которую можно отбросить в окончательном результате.

2) После этого находят значения ϵ , причём необходимо указывать их знаки, как это дано в третьем столбце таблицы, и $\sum \epsilon$. Хотя эта величина не участвует в дальнейших вычислениях, однако $\sum \epsilon$ принято определять. Теоретически алгебраическая сумма отклонений от среднего арифметического должна быть равной нулю, но при вычислении $\sum \epsilon$ это редко наблюдается, так как все ϵ являются числами приближёнными. Однако несомненно, что $\sum \epsilon$ должна быть малой величиной, и если это не имеет места, то можно предполагать наличие арифметической ошибки в предшествующих вычислениях. При определении $\sum \epsilon$ принято находить в отдельности сумму положительных ϵ ($\sum \epsilon^+ = +0,010$) и сумму отрицательных ϵ ($\sum \epsilon^- = -0,013$) и затем определять алгебраическое значение $\sum \epsilon$ в данном примере $\sum \epsilon = -0,003$. Отсюда нетрудно определить сумму абсолютных значений ϵ ($\sum |\epsilon| = 0,023$) и вычислить среднюю ошибку отдельного измерения ($\eta = \pm 0,002$). Величина этой ошибки показывает, что в данном случае значения всех ϵ не выходят за пределы допустимого.

3) Переходя к определению квадратичной или вероятной ошибок результата измерений, вычисляем значения ϵ^2 . При сложении значений ϵ^2 вновь следует выписывать отдельные итоги через каждые пять слагаемых ($39 \cdot 10^{-6}$ и $30 \cdot 10^{-6}$) и затем находить их общую сумму ($69 \cdot 10^{-6}$). Отсюда вычисляем среднюю квадратичную ошибку отдельного измерения ($\sigma = \pm 0,0028$) и квадратичную ошибку

среднего арифметического ($\sigma_R = \pm 0,0009$). Затем выписываем окончательный результат измерений, как это дано в последней строке таблицы. Если бы мы вычислили вероятную ошибку ρ_R , которая равна, очевидно, $\pm 0,0006$, то окончательное значение R мы могли бы дать с четырьмя десятичными знаками, не отбрасывая последней цифры, но мы видим, что она имеет мало значения, так как вероятная ошибка достигает шести единиц в четвёртом десятичном знаке.

Пример 2. При определении абсолютной величины заряда электрона e последовательные измерения некоторой величины N , повторённые двадцать пять раз, дали такие результаты:

61,03	61,19	61,01	61,27	60,97
61,03	61,20	61,07	60,94	60,97
61,16	61,11	61,26	60,98	61,24
60,97	61,05	61,11	61,20	61,24
61,21	61,16	61,39	61,22	61,18

Найдём наиболее вероятное значение N и его вероятную ошибку. Результаты этих вычислений приведены в таблице VI, вполне аналогичной таблице в примере 1.

По отношению к этому вычислению можно сделать следующие замечания:

1) При суммировании значений N можно ограничиться сложением только двух последних цифр и принять во внимание при вычислении $\sum N$, что на месте единиц в пяти случаях стоит нуль (измерения 4, 17, 18, 21 и 22).

2) Сумма положительных и сумма отрицательных ϵ каждая в отдельности имеют большую величину, но алгебраическая сумма ϵ вновь оказывается малой.

3) Определив значение средней ошибки отдельного измерения η и сравнивая её величину со значениями ϵ , мы видим, что ни одно из значений ϵ не имеет ненормально большой величины по сравнению со значением η . Вследствие этого мы можем закончить вычисления, не внося в результаты измерений никаких исправлений.

Таблица VI

N			ε		ε²	
1	61,03	1,40	+0,096	$\sum \epsilon^+ = +1,272$ $\sum \epsilon^- = -1,282$ $\sum \epsilon = -0,010$	0,009216	0,050980
2	61,03		+0,096		09216	
3	61,16		-0,034		01156	
4	60,97		+0,156		24336	
5	61,21		-0,084		07056	
6	61,19	0,71	-0,064		04096	016760
7	61,20		-0,074		05476	
8	61,11		+0,016		00256	
9	61,05		+0,076		05776	
10	61,16		-0,034		01156	
11	61,01	0,84	+0,116		13456	104500
12	61,07		+0,056		03136	
13	61,26		-0,134		17956	
14	61,11		+0,016		00256	
15	61,39		-0,264		69696	
16	61,27	2,61	-0,144		20736	090960
17	60,94		+0,186		34596	
18	60,98		+0,146		21316	
19	61,20		-0,074		05476	
20	61,22		-0,094		08836	
21	60,97	2,60	+0,156		24336	077580
22	60,97		+0,156		24336	
23	61,24		-0,114		12996	
24	61,24		-0,114		12996	
25	61,18		-0,054		02916	
28,16			$\sum \epsilon = 2,560 \quad \eta = \pm 0,102$		$\sum \epsilon^2 = 0,34078$	
$\frac{\sum N}{25} = 61,126$			$\rho_N = \pm \frac{2}{3} \cdot 0,024 = \pm 0,016$		$\sigma = \pm \sqrt{\frac{0,34078}{24}} = \pm 0,119$ $\sigma_N = \pm \frac{0,119}{\sqrt{25}} = \pm 0,024$	
$N = 61,13 \pm 0,02$						

Пример 3. В результате измерений гравитационной постоянной была определена средняя плотность земли D . Критический анализ весьма многочисленных определений этой величины привёл к выводу, что наибольшего доверия заслуживают пять значений D . Эти значения D , выраженные в абсолютной системе единиц, таковы:

$$5,56; 5,527; 5,527; 5,493 \text{ и } 5,507 \text{ см}^{-3} \text{ г.}$$

Найдём на основании этих результатов наиболее вероятное значение D и его вероятную ошибку.

Эти вычисления даны в таблице VII, вполне аналогичной тем, которые мы имели в первых двух примерах.

Таблица VII

D			ε		ε²	
1	5,56	2,614	−0,037	$\sum \epsilon^+ = + 0,046$ $\sum \epsilon^- = - 0,045$ $\sum \epsilon = 0,001$	0,001369	2557
2	5,527		−0,004		0016	
3	5,527		−0,004		0016	
4	5,493		+0,030		0900	
5	5,507		+0,016		0256	
$\frac{\sum D}{5} = 5,5228$			$\sum \epsilon = 0,091 \quad \eta = \pm 0,018$		$\sum \epsilon^2 = 0,002557$	

$$\sigma = \pm \sqrt{\frac{0,002557}{4}} = \pm 0,025,$$

$$\sigma_D = \pm \frac{0,025}{\sqrt{5}} = \pm 0,011, \quad \rho_D = \pm \frac{2}{3} \cdot 0,011 = \pm 0,007$$

$$D = (5,523 \pm 0,007) \text{ см}^{-3} \text{ г}$$

Этот пример представляет интерес в том отношении, что каждое из пяти указанных значений D было получено

в результате больших исследований, в которых отдельные измерения гравитационной постоянной повторялись много раз, и ошибки её окончательных значений были определены. Таким образом, здесь можно применить теорию взвешенных измерений, т. е. можно вычислить наиболее вероятное значение D , установив предварительно вес каждого из пяти его отдельных определений. В результате таких вычислений мы находим значение D , которое отличается от полученного нами, но очень незначительно, менее чем на одну единицу во втором десятичном знаке.

ГЛАВА VI

ОСНОВНЫЕ ПРИЁМЫ ГРАФИЧЕСКОГО АНАЛИЗА РЕЗУЛЬТАТОВ ИЗМЕРЕНИЙ

1. Графическое изображение результатов измерений. Обычные приёмы графического изображения функциональной зависимости между двумя переменными величинами, связанными уравнением

$$y = f(x), \quad (1)$$

очень широко применяются при обработке результатов физических измерений. Измеряя на опыте значения y , отвечающие определённым значениям x , находят соответственные пары значений x_1 и y_1 , x_2 и y_2 , x_3 и y_3 , ..., которые и служат для составления графика. Так, например, можно измерять плотность тела при определённых температурах, или угол отклонения призмой лучей с определённой длиной волны, или радиоактивную силу препарата в определённые моменты времени и т. п. Если результаты этих измерений изобразить графически, то мы получаем кривые, которые очень наглядно представляют изменение: плотности данного тела в определённом интервале температур, или угла отклонения лучей данной призмой в определённом интервале длин волн, или радиоактивной силы данного препарата в течение определённого промежутка времени и т. п. При графическом воспроизведении результатов измерений в громадном большинстве случаев применяется обычная система прямоугольных координат. Другие системы координат, например полярные координаты, также находят применение при графическом изображении результатов измерений, но значительно меньшее.

Более сложными следует считать те случаи, когда закон физического явления выражается формулой, в которую входят три или более переменные величины, т. е. определяется, например, уравнением

$$f(x, y, z) = 0.$$

В этом случае очень часто одной из величин, например z , дают последовательно ряд определённых значений $z_1, z_2, z_3 \dots$ и для каждого из

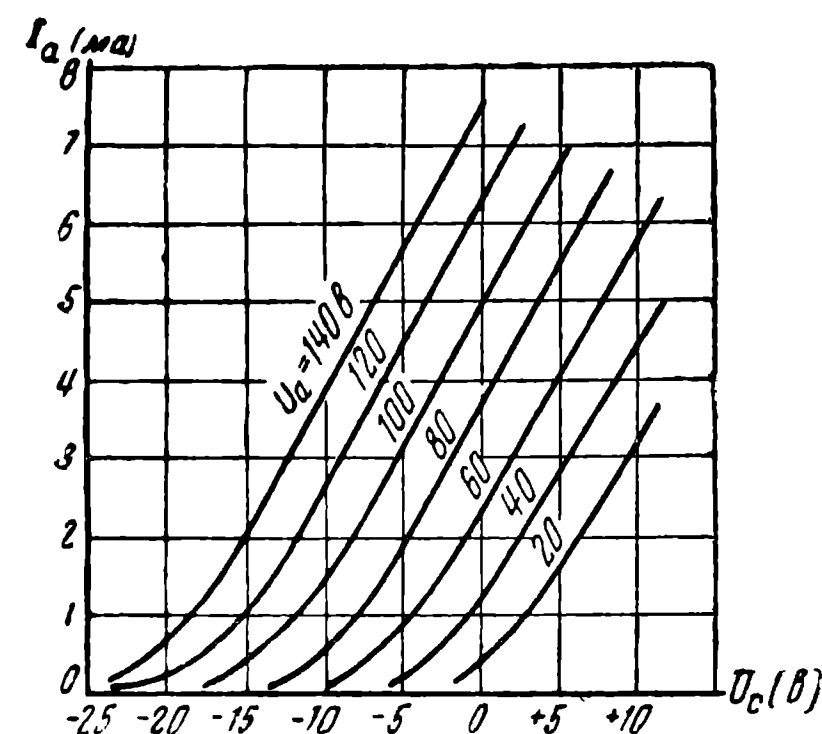


Рис. 12.

силы анодного тока I_a в триодах зависит в основном от двух величин: от напряжения на аноде U_a и от напряжения на сетке U_c . Вследствие этого можно написать:

$$I_a = f(U_a, U_c).$$

При исследовании этих явлений можно, например, анодное напряжение U_a установить на определённую величину и, оставляя его постоянным, измерять силу анодного тока I_a при различных значениях напряжения на сетке U_c . Если результаты этих измерений изобразить графически в прямоугольной системе координат, откладывая на оси абсцисс значения U_c , а на оси ординат значения I_a , то получаются так называемые *сеточные характеристики триода* (рис. 12), каждая из которых отвечает определённому значению анодного напряжения. Таким же приёмом можно исследовать и изобра-

зависимость двух других величин x и y . Если вновь изобразить графически полученные результаты, вычерчивая в одной и той же системе координат кривые для каждого из значений z в отдельности, то получается семейство кривых, каждая из которых изображает зависимость между x и y при определённом значении z . Так, например,

жать графически *анодные характеристики триода*, если, оставляя напряжение на сетке U_c постоянным, измерять силу тока I_a , постепенно изменяя напряжение на аноде U_a .

Результаты измерений, т. е. соответственные значения x и y , наносят на графике в виде точек, которые и служат основой построения кривых. Но, приступая к составлению графика, необходимо иметь в виду, что мы должны заранее исследовать, какие изменения параметров системы можно ожидать при данном физическом явлении. Это объясняется тем, что здесь могут иметь место два основных случая, существенно различных, которые можно разъяснить следующим образом.

1) *Первый случай*. Мы часто встречаем такие явления, при которых течение физических процессов испытывает быстрое, внезапное изменение, что сопровождается такими же быстрыми изменениями тех или иных параметров системы. Можно привести очень много примеров таких явлений:

а) Если мы постепенно изменяем температуру тела, то его объём изменяется, вообще говоря, так же постепенно; но если при изменении температуры в системе возникают фазовые превращения, т. е. если твёрдое тело при нагревании плавится или жидкость при охлаждении кристаллизуется, то объём системы обнаруживает резкое изменение, хотя её температура остаётся при этом постоянной.

б) Если мы постепенно повышаем напряжение на обкладках конденсатора, то его заряд возрастает также постепенно; но если в системе возникает искровой разряд, то её электрическое состояние испытывает резкое изменение.

в) Явления радиоактивности объясняются особыми процессами, которые протекают внутри атомов радиоактивных тел, так называемым *радиоактивным распадом атомов*; исследования показали, что все свойства каждого атома при его распаде испытывают резкое, внезапное изменение.

Таким образом, мы видим, что фазовые превращения, искровой разряд, радиоактивный распад атомов и многое другое могут служить примерами тех физических явлений, при которых параметры системы испытывают внезапные

изменения. Очевидно, что во всех подобных случаях на соответствующих графиках должны наблюдаться резкие изгибы кривых, резко выраженные максимумы или минимумы. Отсюда следует, что, нанеся экспериментальные точки на графике, необходимо очень внимательно рассмотреть их относительное расположение. Если точки сильно разбросаны, хотя бы в одной части графика, и есть основание предполагать, что в данном случае могут иметь место особые явления, аналогичные указанным выше, то такое расположение точек подтверждает наше предположение. Если затем построить возможно более точный график, то мы можем эти особые явления исследовать количественно, хотя очень часто такое исследование оказывается недостаточно точным и в силу этого может носить только предварительный характер.

2) В т о р о й с л у ч а й. Однако очень часто мы имеем все основания считать, что данное физическое явление имеет вполне плавное течение, при котором все параметры системы изменяются постепенно. Так, объём или плотность тела при постепенном его нагревании, пока в системе не наблюдается фазовых превращений, изменяются также постепенно; справедливость этого вывода подтверждается непосредственными измерениями. Точно так же сила анодного тока в триоде при постоянном напряжении на аноде и постепенном изменении напряжения на сетке изменяется также постепенно, что подтверждается соответствующими измерениями и т. п. В подобных случаях экспериментальные точки, нанесённые на графике, должны были бы ложиться на некоторую плавную кривую, т. е. на кривую, не имеющую резких искривлений, которая и отвечает данному процессу. Однако практика показывает, что некоторый разброс экспериментальных точек на графиках всегда наблюдается. Вследствие этого, если бы мы соединяли непрерывной линией все экспериментальные точки на графике, то получались бы кривые очень неправильного вида с резкими перегибами и искривлениями, как это изображено на рис. 13 пунктирной линией. Так как мы рассматриваем те процессы, которые носят плавный характер, то разброс точек на графике и соответствующие искривления на кривой мы можем объяснить только

ошибками наших измерений. Справедливость этого заключения подтверждается в весьма разнообразных случаях тем, что, если мы повторяем те же измерения более точным методом, так что их ошибки уменьшаются, то уменьшается и разброс точек на графике. Таким образом, мы можем считать, что если бы мы имели результаты измерений, совершенно свободные от ошибок, то точки, отвечающие этим результатам,

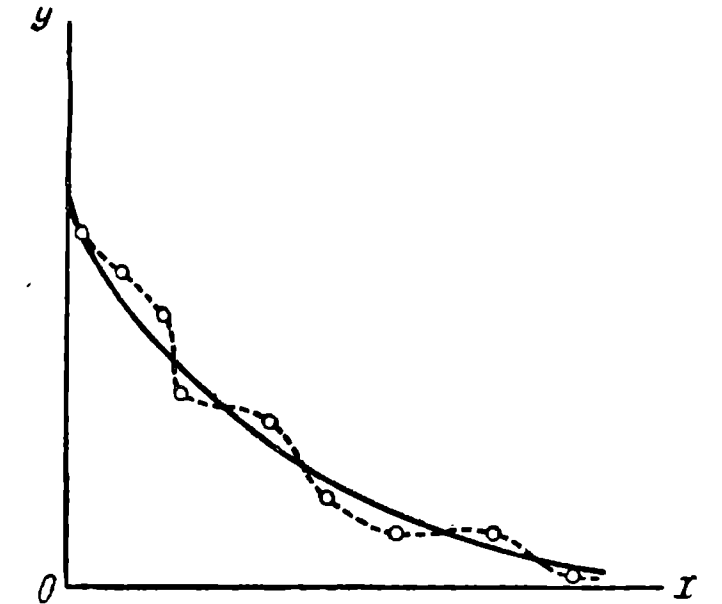


Рис. 13.

строго ложились бы на соответствующую плавную кривую. Вследствие этого все процессы, которые имеют заведомо плавное течение, принято изображать графически также плавными кривыми, проводя их не через экспериментальные точки, а так, чтобы кривая, как обычно говорят, проходила возможно ближе ко всем точкам на графике, как это показано на рис. 13 сплошной линией. Этот приём, которому при проведении плавных кривых необходимо всегда следовать, мы рассмотрим более подробно; для этого воспользуемся следующими соображениями.

Каждая из точек, нанесённых на графике, отвечает определённым значениям x и y . Эти значения мы получаем в результате измерений; вследствие этого они содержат некоторые ошибки. Ошибку имеет не только y , но и x , так как, хотя мы и говорим, что находим значения y при определённых значениях x , но эти значения x мы получаем также из измерений, т. е. приближённо. Так, если мы измеряем плотность тела при определённых температурах, то последние мы отсчитываем по термометру, т. е. с некоторой ошибкой; точно так же, если мы определяем радиоактивную силу препарата как функцию времени, то промежутки времени мы измеряем по секундомеру или часам, т. е. тоже с некоторой ошибкой и т. п. Вследствие этого при построении графика мы должны считаться с ошибками как y , так и x . В ре-

зультате мы должны наносить на графике не точки (x_1, y_1) , (x_2, y_2) , (x_3, y_3) , ..., а небольшие прямоугольники (рис. 14, а), стороны которых равны $(2\epsilon_x, 2\epsilon_y)$, $(2\epsilon'_x, 2\epsilon'_y)$, $(2\epsilon''_x, 2\epsilon''_y)$, ..., т. е. равны удвоенным абсолютным ошибкам x и y при их последовательных измерениях, например, квадратичными или вероятными ошибками x и y . В таком случае можно утверждать, что те точки (x, y) , которые мы могли бы получить при измерениях, свободных от ошибок, должны находиться где-то внутри этих прямоугольников.

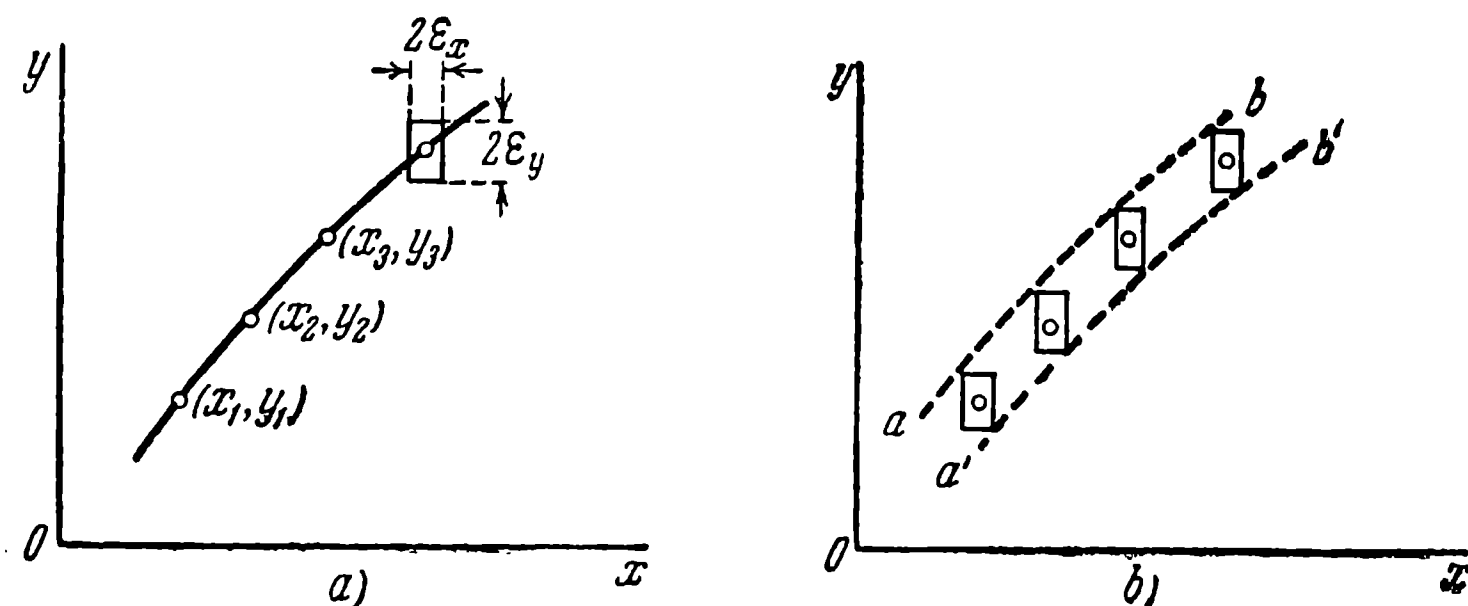


Рис. 14.

Иначе говоря, если мы проведём две предельные кривые ab и $a'b'$ (рис. 14, б), то кривая, точно отвечающая виду функции f , может не проходить через экспериментальные точки, однако она не должна выходить за пределы узкой полосы, ограниченной кривыми ab и $a'b'$. Отсюда следует, что теория ошибок даёт нам право проводить плавные кривые, как было указано выше, но возможные отклонения кривой от экспериментальных точек, вообще говоря, должны зависеть от величины ошибок x и y . На практике обыкновенно ограничиваются нанесением точек, отвечающих экспериментальным данным, и затем строят плавную кривую без резких искривлений. Вначале кривую вычерчивают карандашом в виде достаточно широкой полосы, которая постепенно выравнивается при помощи подчистки и исправлений. Затем кривая окончательно вычерчивается по частям при помощи лекал, карандашом или тушью.

При построении графика иногда оказывается, что одна или две пары значений x и y дают точку, сильно удалённую от кривой, т. е. некоторые значения y , полученные из измерений при определённых значениях x , вызывают отклонения точки от кривой на величины, значительно превышающие возможную ошибку измерений. Если нет оснований предполагать, что в данной области может иметь место резкое изменение в ходе исследуемой зависимости y от x (стр. 203), то такие отклонения можно объяснить только какой-либо грубой ошибкой в измерениях, в отсчётах или в вычислениях. В этих случаях следует, прежде всего, проверить вычисления. Если в них ошибок не оказывается, а повторить измерения не представляется возможным, то приходится такую точку считать ошибочной и совершенно не принимать её во внимание (стр. 194). Если же измерения для этой точки можно повторить ещё раз, то их необходимо произвести особенно тщательно, что обыкновенно даёт возможность установить ошибку и внести необходимые исправления. Но если при повторных измерениях будет всё же получен прежний результат, то к нему необходимо отнестись очень внимательно, так как такие результаты могут указывать на те резкие нарушения в ходе процесса, о которых мы говорили выше и наличия которых в данном случае мы не предполагали. В этих случаях необходимо исследовать весь ближайший участок кривой, т. е. несколько изменять в ту и другую сторону значения x и определять из измерений соответствующие значения y .

Графические изображения результатов измерений не только дают очень наглядное представление о взаимной зависимости исследуемых величин, но одновременно могут служить и для измерительных целей, так как по графику можно находить значения одной величины, если значения остальных величин известны. Графики, которые служат измерительным целям, необходимо вычерчивать с большой точностью и строить в достаточно большом масштабе.

В дальнейшем будем иметь в виду преимущественно графики последнего типа, т. е. точные графики для отсчётов.

При графическом изображении результатов измерений необходимо иметь в виду следующее:

1) Чрезвычайным облегчением при построении графиков является применение миллиметровой бумаги, а также других видов координатных бумаг, например логарифмической полярной и т. д. (стр. 231), но при этом необходимо принимать во внимание возможные неточности сетки, нанесённой на бумаге. Так, сетка на миллиметровой бумаге, которая особенно широко применяется при построении графиков, далеко не всегда отвечает нормальной длине сантиметра, и часто имеют место случаи, когда длина делений на бумаге в продольном и поперечном направлениях оказывается различной. Вследствие этого рекомендуется миллиметровую бумагу предварительно исследовать, пользуясь чертёжным циркулем и точной миллиметровой линейкой. Так как при построении графиков на миллиметровой бумаге мы наносим точки по её делениям, то и отсчёты расстояний на графиках следует делать по делениям на бумаге, а не при помощи циркуля и линейки, хотя бы бумага и была предварительно исследована. Ошибка при отсчёте по графикам, вычерченным на миллиметровой бумаге, может достигать $\pm 0,25$ мм, а часто бывает и больше этой величины. Если же график точно рассчитан и тщательно вычерчен на листе плотной белой бумаги, то ошибку отсчёта можно уменьшить до $\pm 0,1$ мм.

2) При построении графика весьма большое значение имеет удачный выбор масштаба. Вообще говоря, масштаб на графике по той и другой оси мы можем выбирать как угодно, и тем не менее при выборе масштаба необходимо иметь в виду определённые условия:

а) Если нам необходимо, чтобы точность отсчётов по графику отвечала точности наших измерений, то наименьшее расстояние, которое мы можем отсчитывать на графике, должно соответствовать величине абсолютной ошибки наших измерений. Так, например, если ошибка отсчётов на графике достигает $\pm 0,25$ мм, то абсолютную ошибку наших измерений и следует принять на графике равной этой величине. Отсюда и определяется соответствующий масштаб. Однако это требование, которым, казалось бы,

следовало всегда руководствоваться, далеко не во всех случаях можно выполнить на практике и часто приходится от него отступать, обычно в сторону уменьшения масштаба. Отсюда следует, что отсчёты по графику могут быть менее точными, чем измерения, для которых он составлен.

б) Не следует строить графики в излишне большом масштабе. Листы миллиметровой бумаги с размерами, обычными для настольного блокнота, т. е. 20 см на 30 см, в большинстве случаев являются вполне достаточными.

в) Относительную величину масштаба на той и другой оси следует выбирать так, чтобы кривые не были слишком плоскими, растянутыми вдоль одной из осей. Относительная ошибка отсчётов по таким графикам может заметно возрастать, кроме того, они оказываются менее наглядными. Если, например, кривая очень вытянута вдоль оси y и круто падает к оси абсцисс (рис. 15), то незначительная неточность при отсчёте x или недостаточно точно взятое направление ординаты может дать очень большую ошибку при измерении y . Точно так же из двух кривых (рис. 16), которые представляют зависимость вторичной электронной эмиссии от скорости первичных электронов, нижняя кривая даёт менее наглядное изображение этого явления из-за неудачного выбора относительной величины масштаба на осях.

3) При построении графика, вообще говоря, нет оснований брать особенно большое число точек, так как ход плавной кривой в большинстве случаев выявляется достаточно определённо, если на протяжении всего интервала изменений x мы возьмём для него десять-пятнадцать точек. Но если на кривой намечается возможность таких точек как максимум, минимум, точки перегиба, то точки на графике следует наносить значительно чаще, так как иначе ход кривой в таких особых точках утрачивает определённость.

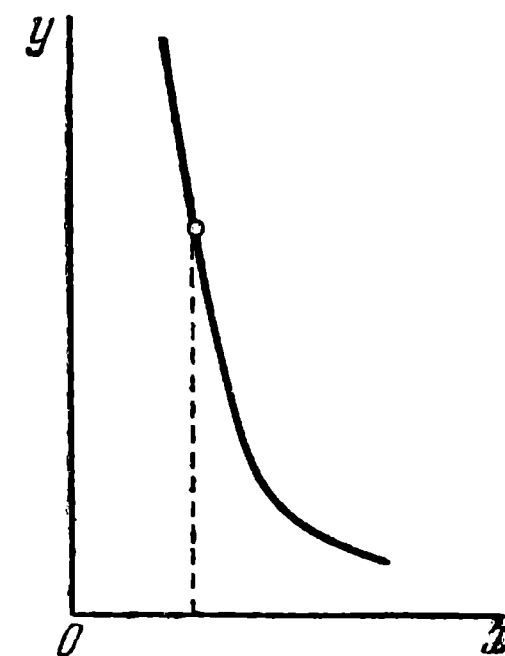


Рис. 15.

4) Наконец, необходимо указать ещё одну особенность графиков.

Для этого рассмотрим, например, график изотермического процесса в газах; при температурах значительно выше критической этот процесс с достаточной точностью может быть представлен законом Бойля—Мариотта:

$$p \cdot v = a,$$

где p и v обозначают давление и объём газа, а a обозначает некоторый постоянный коэффициент. С математиче-

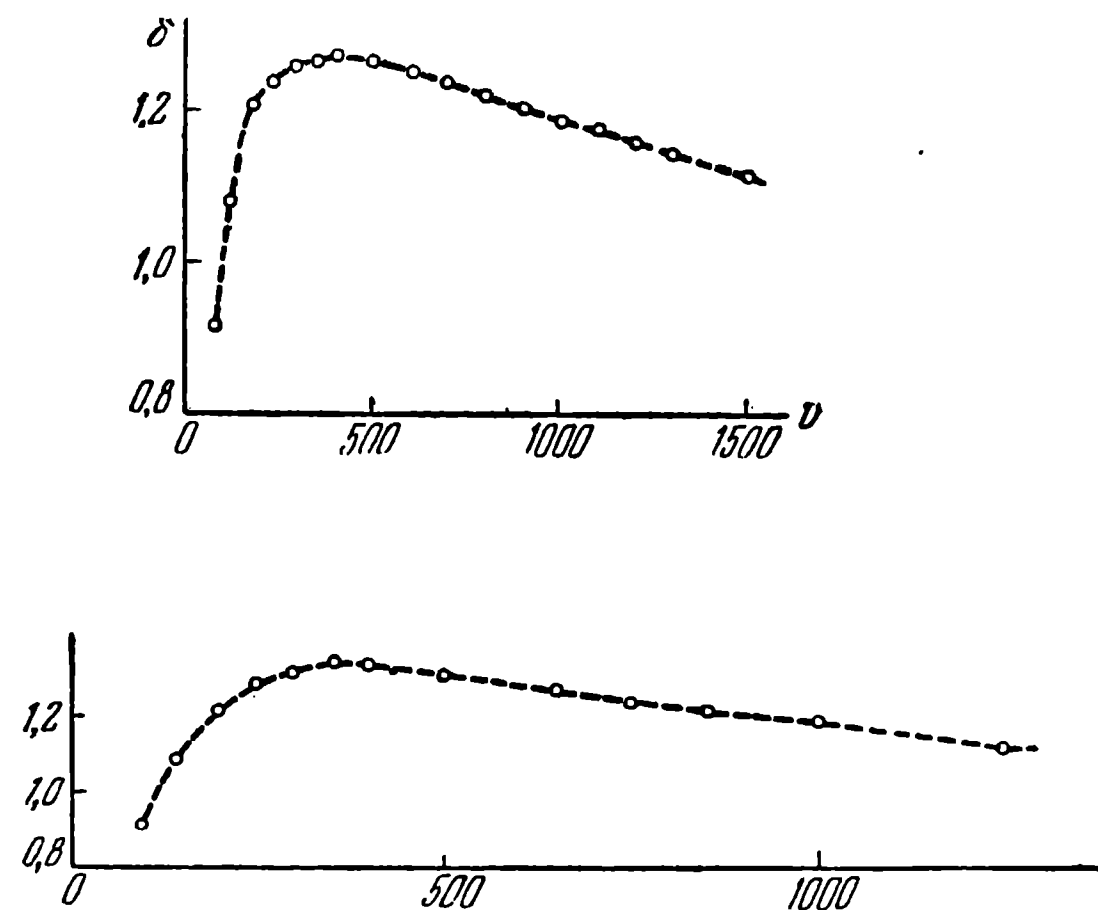


Рис. 16.

ской стороны это уравнение можно рассматривать как частный случай уравнения, которым определяется результат умножения двух чисел p и v .

Если это уравнение изобразить графически в прямоугольной системе координат, то при одинаковом масштабе на обеих осях получаются, как известно, равносторонние гиперболы, асимптотами которых служат оси координат (рис. 17). Для того чтобы выполнить построение этих гипербол, необходимо знать соответствующие числовые значения постоянного параметра a . Так,

средняя гипербола mn на рис. 71 построена в предположении, что значение a равно пяти.

График, который получается в результате этого построения, мы можем применять для определения значений p при данных значениях v и обратно, но только при том частном значении a , для которого он построен. Отсюда следует, что если постоян-

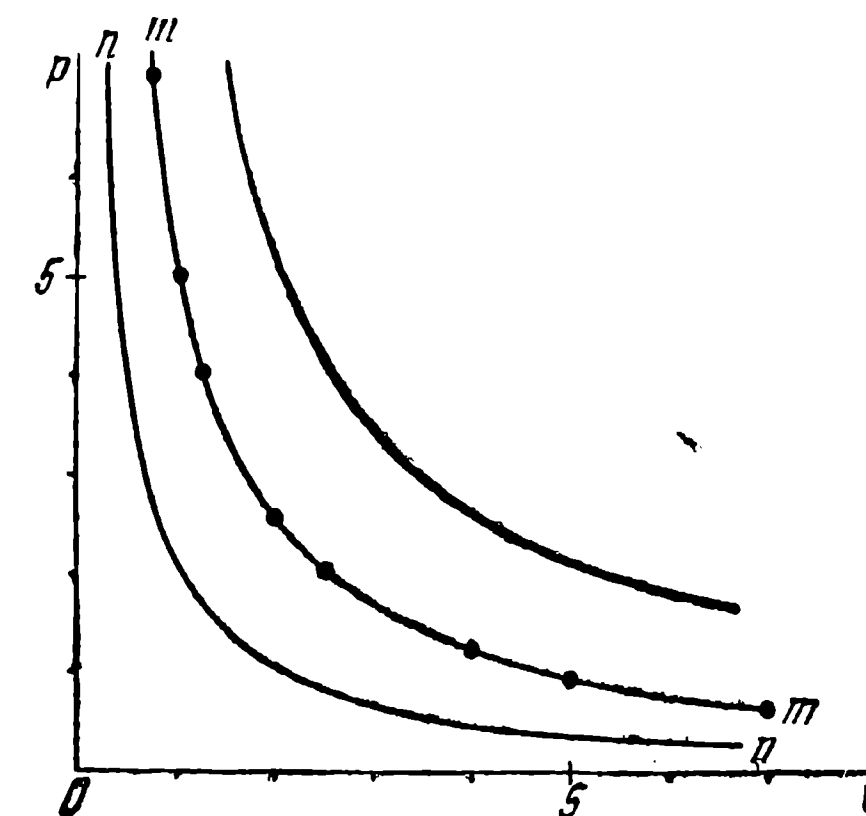


Рис. 17.

ная a получает иное значение, то необходимо построить новый график, в данном случае новую гиперболу; так, например, нижняя гипербола nn на рис. 17 отвечает значению a , равному двум.

Таким образом, мы видим, что каждый график отвечает определённым числовым значениям постоянных коэффициентов в соответствующем уравнении, иначе говоря, графики можно применять для измерений только при тех значениях коэффициентов, для которых данный график был построен.

2. Графическое дифференцирование и интегрирование.

Точно выполненное графическое изображение результатов измерений даёт возможность производить ряд измерительных и вычислительных операций. Эти операции, как

правило, можно выполнять лишь с относительно небольшой точностью, тем не менее графические измерения и вычисления очень широко применяются на практике, так как они являются весьма простыми, хотя и требуют обычно вспомогательных построений.

Так, по графику мы можем определять значения y при любых значениях x и наоборот, если эти значения переменных лежат в пределах, для которых составлен график. Точно так же по графику мы можем измерять максимальные и минимальные значения той или другой величины, а также все те значения одной величины, которые соответствуют максимуму или минимуму другой величины, хотя бы таких измерений непосредственно и не было произведено. При графических измерениях рекомендуется, кроме циркуля и чертёжной линейки, иметь ещё прозрачную (целлулоид или плексигласс) линейку с миллиметровой шкалой и прозрачную пластинку с координатными осями или лучше с координатной сеткой.

Чрезвычайно важным является, далее, то обстоятельство, что, пользуясь чисто графическими приёмами, мы можем производить такие вычислительные операции, как дифференцирование и интегрирование. Эти вопросы требуют более подробных разъяснений, которые в элементарной форме можно изложить следующим образом.

Если аналитическое выражение функции

$$y = f(x) \quad (1)$$

не известно, и мы имеем только точный график этой функции в прямоугольной системе координат, то для решения общей задачи о нахождении функций

$$\varphi(x) = \frac{df(x)}{dx} \quad (2)$$

и

$$F(x) = \int f(x) dx \quad (3)$$

мы можем воспользоваться только графическими приёмами.

Но, кроме общей задачи о построении кривых (2) и (3), могут встречаться и более частные вопросы, во-первых, о графическом вычислении производной (2) при определённом значении x и, во-вторых, о графическом вычислении интеграла (3), взятого в определённых пределах, т. е. о графическом вычислении площади, ограниченной кривой (1), осью абсцисс и двумя определёнными ординатами.

Графическое решение этих частных вопросов оказывается более простым, и вместе с тем при решении этих вопросов мы получаем общие указания на те графические построения, которые приводят к решению и общих задач о построении кривых (2) и (3). В такой последовательности мы и рассмотрим вопросы о графическом дифференцировании и интегрировании.

I. *Графическое дифференцирование.* Для графического определения значения производной функции (1) при определённом значении x_1 (рис. 18) следует построить

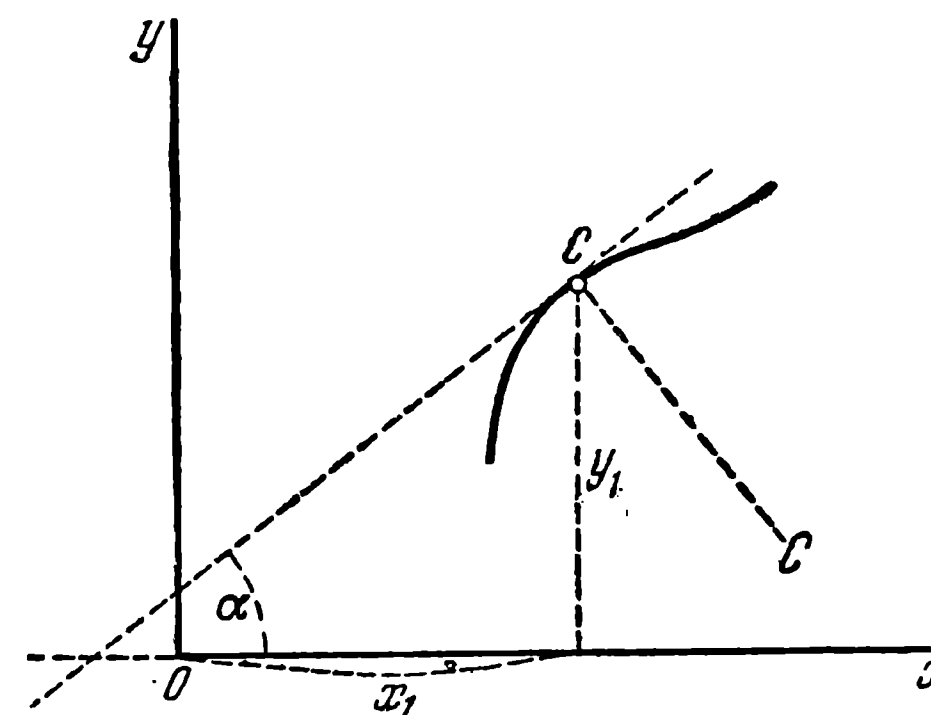


Рис. 18.

касательную к кривой в точке (x_1, y_1) и определить угол α , образуемый этой касательной с положительным направлением оси x . Это следует из того, что, как известно,

$$\frac{df(x)}{dx} = \operatorname{tg} \alpha. \quad (4)$$

Эта задача, весьма простая в теории, графически может быть решена, вообще говоря, лишь очень приближённо,

так как провести касательную в данной точке произвольной кривой можно лишь с очень ограниченной точностью. При таких построениях можно рекомендовать предварительно наметать нормаль CC , считая небольшой участок кривой вблизи точки (x_1, y_1) дугой конического сечения, в простейших случаях — дугой окружности. Построение нормали действительно можно выполнить несколько точнее, но и при этом условии провести касатель-

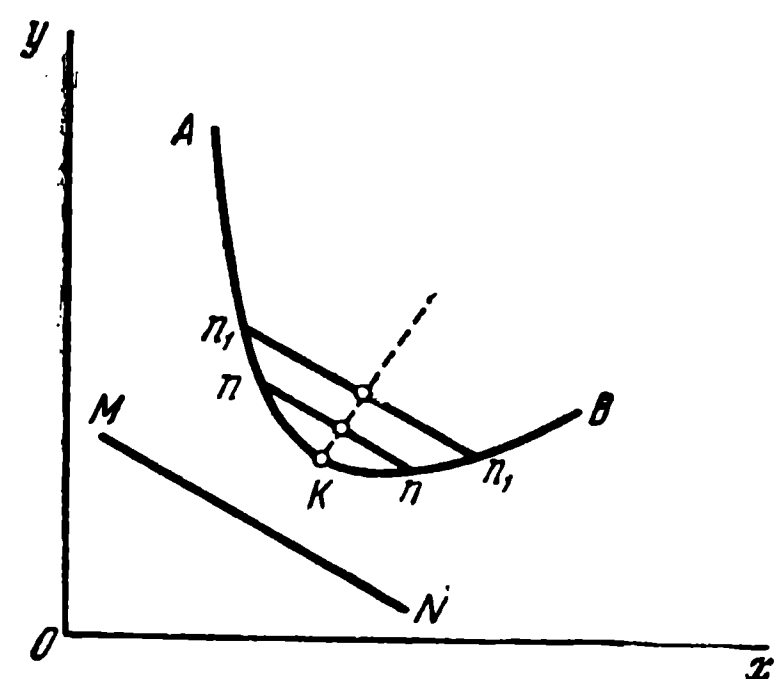


Рис. 19.

ную удаётся всё же лишь очень приближённо. Вследствие этого, если при таких задачах и применяются графические решения, то их результаты приходится расценивать как неточные и пользоваться ими следует только при предварительных расчётах.

Одновременно та же задача графического дифференцирования ставит на очередь вопрос о построении касательных, парал-

лельных заданным направлениям, и об определении их точек касания. Для решения этого вопроса пользуются вспомогательным построением, которое заключается в следующем.

Если дана кривая AB (рис. 19) и то направление MN , параллельно которому необходимо провести касательную, то следует построить две близкие секущие (хорды) nn_1 и n_1n_2 , параллельные MN , и разделить их пополам. Считая, что часть кривой n_1kn_2 мало отличается от дуги конического сечения, мы можем воспользоваться теоремой о сопряжённых диаметрах. Таким образом, если соединить середины хорд прямой, то её пересечение с кривой ab даёт точку касания касательной, параллельной MN .

Этот приём можно применять для построения дифференциальных кривых. Если кривая AB (рис. 20) отвечает функции (1) и мы должны построить её дифференциальную кривую (2), то следует, прежде всего, наметить направ-

ления для ряда касательных в различных точках кривой AB и, пользуясь указанным приёмом, найти точки их касания $C_1, C_2, C_3, \dots$, причём самих касательных можно и не проводить; из точек $C_1, C_2, C_3, \dots$ следует провести прямые, параллельные оси y . Затем из некоторой точки, расположенной на продолжении оси x , или из так называемого полюса P проводим ряд лучей $1, 2, 3, \dots$, параллельных выбранным направлениям для касательных; лучи

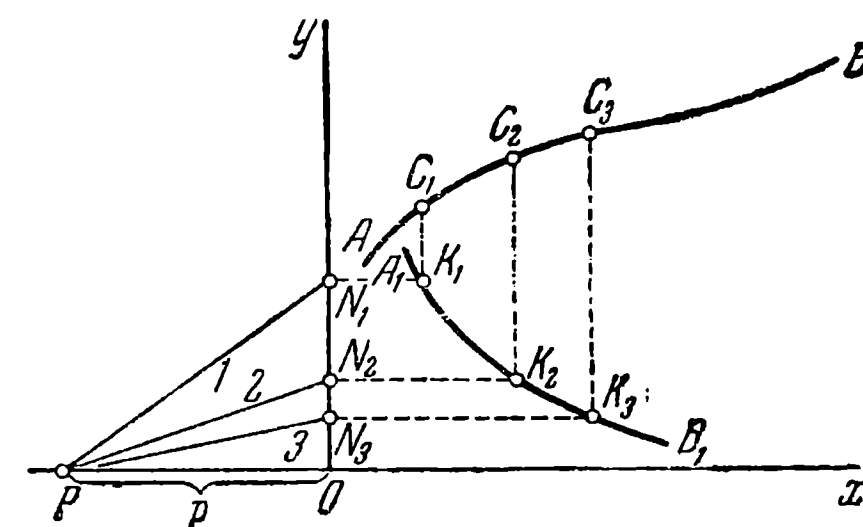


Рис. 20.

продолжаем до пересечения с осью y . Так как все лучи образуют с осью x те же углы, как и отвечающие им касательные, то мы можем применить к лучам уравнение (4), т. е. можно утверждать, что лучи пересекают ось y в точках $N_1, N_2, N_3, \dots$, расстояния которых от начала координат должны быть пропорциональными тангенсам углов α . Вследствие этого, проводя из точек $N_1, N_2, N_3, \dots$ прямые, параллельные оси x , мы получаем в пересечении их с соответствующими вертикальными прямыми точки $K_1, K_2, K_3, \dots$ дифференциальной кривой, которую и строим по этим точкам в виде плавной кривой A_1B_1 . Масштаб этой кривой по оси x остаётся одинаковым с масштабом кривой AB по той же оси, что же касается масштаба дифференциальной кривой по оси y , то этот масштаб ∂m_y определяется из следующего равенства, которое мы приводим без вывода:

$$\partial m_y = \frac{m_x \cdot m_y}{p}, \quad (5)$$

где m_x и m_y обозначают масштаб соответственно по оси x

и по оси y , выбранный для кривой AB , а p обозначает полюсное расстояние OP . Все величины в формуле (5) принято выражать в миллиметрах.

Такое построение не даёт точных дифференциальных кривых, и даже при удачном выборе точек касания C и при некотором навыке в таких построениях мы всё же обычно получаем недостаточно точные результаты.

II. *Графическое интегрирование.* Графическое определение площади, заключённой между осью абсцисс, частью кривой AB (рис. 21, а) и двумя ординатами y_1 и y_2 , производится при помощи сложения площадей узких поло-

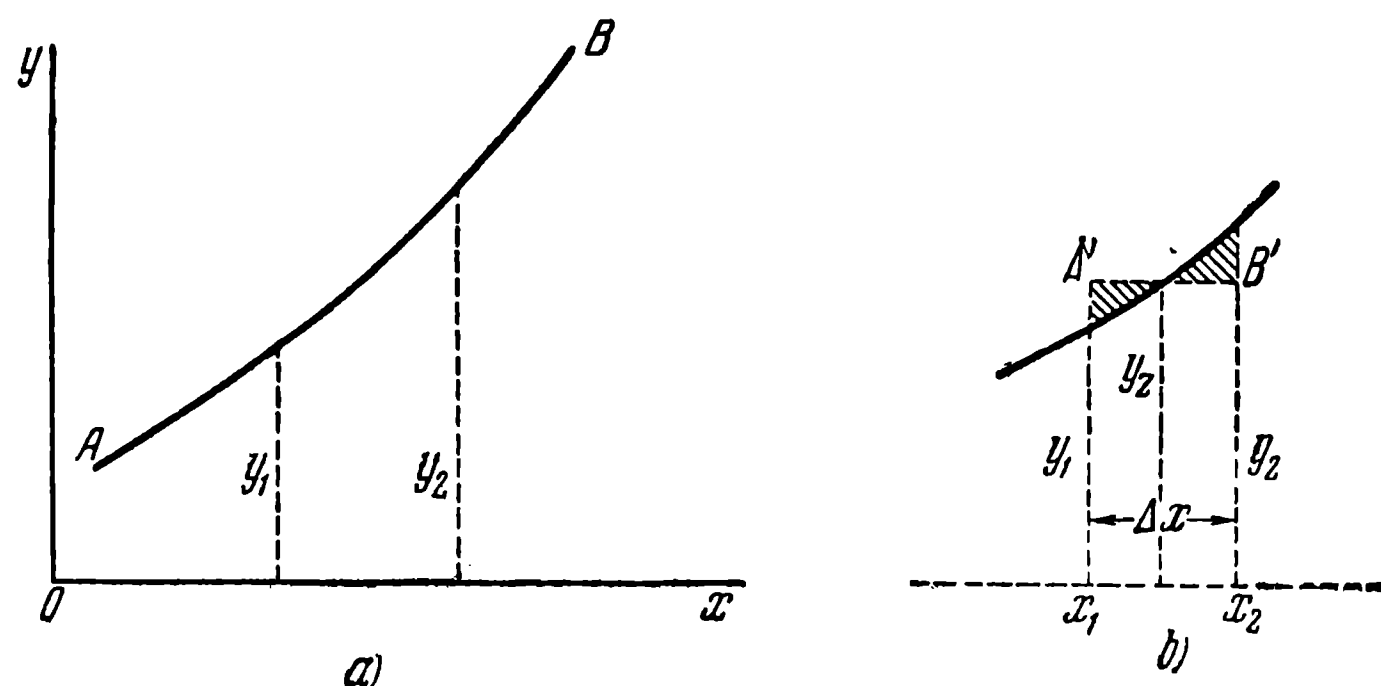


Рис. 21.

сок, на которые разбивается рядом ординат вся определяемая площадь; ширина этих полосок может быть различной. Вверху каждой полоски проводят параллельно оси x отрезок $A'B'$ так, чтобы заштрихованные площади (рис. 21, б) были равны. При этом условии площадь каждой из полосок можно считать равной площади соответствующего прямоугольника, основание которого равно ширине полоски Δx , а высота равна ординате точки пересечения с отрезком $A'B'$. Общая сумма площадей этих прямоугольников и даёт величину определяемой площади.

Равенство заштрихованных площадей очень часто удается установить с достаточной точностью непосредственно на глаз, если основания прямоугольников являются

достаточно малыми и удачно выбраны. Но в некоторых случаях непосредственно оценить равенство треугольных площадей, которые лежат по ту и другую сторону кривой, было бы затруднительным. Вследствие этого для более точного определения площади S отдельных полосок применяется формула, в основе которой лежит предположение, что часть кривой в границах ширины полосок можно считать отрезком кубической параболы. Эта формула, которую мы оставляем без вывода, имеет такой вид:

$$S = \frac{\Delta x}{3} (y_m + 2y_z), \quad (6)$$

где y_m обозначает полусумму крайних ординат полоски, т.е.

$$y_m = \frac{y_1 + y_2}{2},$$

а y_z обозначает ординату средней точки полоски, т.е. ординату, отвечающую абсциссе:

$$\left(x_1 + \frac{\Delta x}{2}\right).$$

Формулу (6) можно считать практически точной, если величины Δx взяты достаточно малыми.

Такая замена кривой (1) ступенчатой линией, которая состоит из горизонтальных и вертикальных отрезков, даёт возможность графически определять интегральную кривую (3). Этот вопрос мы изложим очень коротко, ограничиваясь чисто практическими указаниями.

Кривую AB (рис. 22) мы заменяем ступенчатой линией, которую строим совершенно тем же приёмом, но величины Δx в этом случае можно брать несколько большими. Вначале следует построить интегральную кривую для этой ступенчатой линии. Для этого продолжаем её ступени в вертикальном и горизонтальном направлениях; последние отрезки проводим до пересечения с осью y . Так как производная интегральной функции должна равняться интегрируемой функции, то интегральная функция в промежутке между x_1 и x_2 должна изображаться отрезком прямой, причём угол α_1 наклона этой прямой к оси x должен удовлетво-

определяются, очевидно, видом функции f . Следует обратить внимание на то, что на функциональной шкале откладываются значения функции y , а деления шкалы обозначаются значениями аргумента x .

Для того чтобы построить графически функциональную шкалу, необходимо выбрать определённую единицу для измерения функции, иначе говоря, необходимо длину некоторого отрезка прямой принять за единицу для изме-



Рис. 23.

рения y . Эта величина, т. е. длина отрезка, принятого за единицу для измерения y , называется модулем шкалы. Модуль шкалы обыкновенно выражается в миллиметрах; так, если модуль шкалы равен 15,0 мм и мы должны построить значения y , равные, например, 1,0; 2,1; 3,4 и 5,5, то нам необходимо отложить на шкале от точки A отрезки длиной l , равной:

$$\begin{aligned} l_1 &= 15 \text{ мм} \cdot 1,0 = 15,0 \text{ мм}, & l_3 &= 15 \text{ мм} \cdot 3,4 = 51,0 \text{ мм}, \\ l_2 &= 15 \text{ мм} \cdot 2,1 = 31,5 \text{ мм}, & l_4 &= 15 \text{ мм} \cdot 5,5 = 82,5 \text{ мм}. \end{aligned}$$

Длину отрезка функциональной шкалы, отвечающего числовому значению функции, равному y , принято обозначать $\bar{y}$ или $\overline{f(x)}$. Таким образом, обозначая модуль шкалы через m , мы можем, очевидно, написать:

$$\bar{y} = my \text{ или } \overline{f(x)} = my = mf(x). \quad (8)$$

Это равенство называют уравнением шкалы; оно служит для определения модуля шкалы, величину которого всегда необходимо выбирать так, чтобы общая длина всей шкалы не превышала определённых размеров. Это можно сделать на основании следующих соображений: пусть крайние значения аргумента равны x_0 и x_n ; соответственно этому крайние значения y определяются выражениями

$$y_0 = f(x_0) \text{ и } y_n = f(x_n).$$

Если модуль шкалы обозначим попрежнему через m , то расстояние между крайними точками на шкале, отвечающими значениями y_0 и y_n , согласно уравнению (8) равно:

$$\bar{y}_n - \bar{y}_0 = mf(x_n) - mf(x_0).$$

Это расстояние определяет длину всей шкалы; если она не должна превышать N миллиметров, то можно написать:

$$m[f(x_n) - f(x_0)] \leq N.$$

Отсюда находим значение модуля шкалы:

$$m \leq \frac{N}{f(x_n) - f(x_0)} \leq \frac{N}{y_n - y_0}, \quad (9)$$

т. е. модуль шкалы не должен быть больше её общей длины, разделённой на разность крайних значений функции.

Применяя уравнение (9) для определения модуля шкалы, следует общую длину шкалы N выражать в миллиметрах, так как эта единица принята для измерения модуля. Что же касается разности в знаменателях правой части выражения (9), то при анализе физических явлений функция y обыкновенно обозначает какую-либо физическую величину, например длину, силу, электрическое сопротивление и т. п., и как физическая величина имеет определённую размерность. Однако в данном случае мы оперируем с функцией y как с математической величиной, и в знаменатель правой части уравнения (9) следует подставлять только числовое значение разности $(y_n - y_0)$, не принимая во внимание размерности физической величины, определяемой функцией y .

Найдём, например, модуль шкалы функции, определяющей закон свободного падения тела без начальной скорости:

$$s = \frac{gt^2}{2}, \quad (10)$$

для значений t в интервале от 0 до 10 сек.; допустим при этом, что общая длина шкалы не должна превышать 20 см. Принимая g равным $9,8 \text{ м сек}^{-2}$, находим крайние

значения s , отвечающие значениям t , равным 0 и 10 секундам:

$$s_0 = 0 \text{ м} \text{ и } s_{10} = \frac{9,8 \text{ м сек}^{-2} \cdot 10^2 \text{ сек}^2}{2} = 490 \text{ м}.$$

Применяя формулу (9), находим значение модуля m :

$$m \leq \frac{200 \text{ мм}}{490} = 0,41 \text{ мм},$$

т. е. m оказывается несколько больше 0,4 мм. Таким образом, если модуль шкалы в данном случае взять равным 0,4 мм, иначе говоря, если на шкале функции (10) метр пути принять равным 0,4 мм, то общая длина шкалы окажется несколько меньшей 20 см.

Для расчёта и построения функциональных шкал можно применять два различных приёма.

1) Если аналитическое выражение функции $f(x)$ известно, то построить её функциональную шкалу можно тем приёмом, который был указан выше, т. е. мы даём x в интервале его изменений ряд равноотстоящих значений, вычисляем соответствующие значения y и, выбрав определённый модуль шкалы, строим их на некоторой прямой. Так, в примере свободного падения тела без начальной скорости для построения шкалы функции (10) в указанном выше интервале времени даём t ряд последовательных значений, равных 0, 1, 2, ..., 10 секунд, и вычисляем соответствующие значения s . Получаем результаты, которые приведены в таблице VIII.

Таблица VIII

$t \text{ сек}$	0	1	2	3	4	5	6	7	8	9	10
$s \text{ м}$	0	4,9	19,6	44,1	78,4	122,5	176,4	240,1	313,6	396,9	490,0

Принимая модуль шкалы равным попрежнему 0,4 мм, вычисляем длину отрезков, отвечающих всем этим значениям s , и откладываем их на прямой AB (рис. 24); деления шкалы отмечаем значениями t .

Точно так же мы можем поступить в том случае, если строим шкалу какой-либо тригонометрической функции в определённом интервале значений её аргумента, например функции

$$y = \cos \alpha.$$

В этом случае, ограничиваясь первой четвертью окружности, мы можем значения y , например через каждые 5° ,

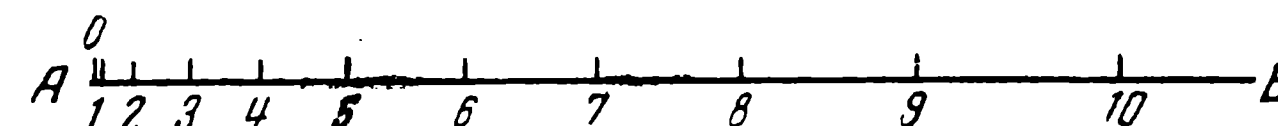


Рис. 24.

взять непосредственно из таблиц натуральных косинусов и построить соответствующую функциональную шкалу. Модуль этой шкалы, очевидно, равен общей длине шкалы, так как предельное значение косинуса равно единице, и мы строим шкалу в первой четверти окружности. Шкала косинусов дана на рис. 25, её модуль принят равным 100 мм.



Рис. 25.

2) В последнем примере мы можем поступить и иначе, а именно, возьмём первый квадрант некоторой окружности, радиус которой примем равным единице (рис. 26), и отложим на ней от точки B значения α через каждые 10° ; получаем точки деления окружности, отвечающие углам $0^\circ, 10^\circ, 20^\circ, \dots, 90^\circ$. Если взять проекции этих точек на горизонтальную прямую MN , то их расстояния до центра окружности O дают значения косинусов углов, равных соответственно $0^\circ, 10^\circ, 20^\circ, \dots, 90^\circ$. Таким образом, шкала MN , деления которой мы обозначаем значениями α , является шкалой функции $\cos \alpha$ в первой четверти окружности; эту шкалу мы можем расположить, конечно, так, чтобы её деления возрастали в направлении слева направо, как на обычных шкалах.

Этот способ построения функциональных шкал можно назвать графическим. В тех случаях, когда аналитическое выражение функции y остаётся неизвестным и дано только её графическое изображение, этот способ является очевидно, единственно возможным. Такие условия мы встречаем, если, измеряя какую-либо величину (функцию)

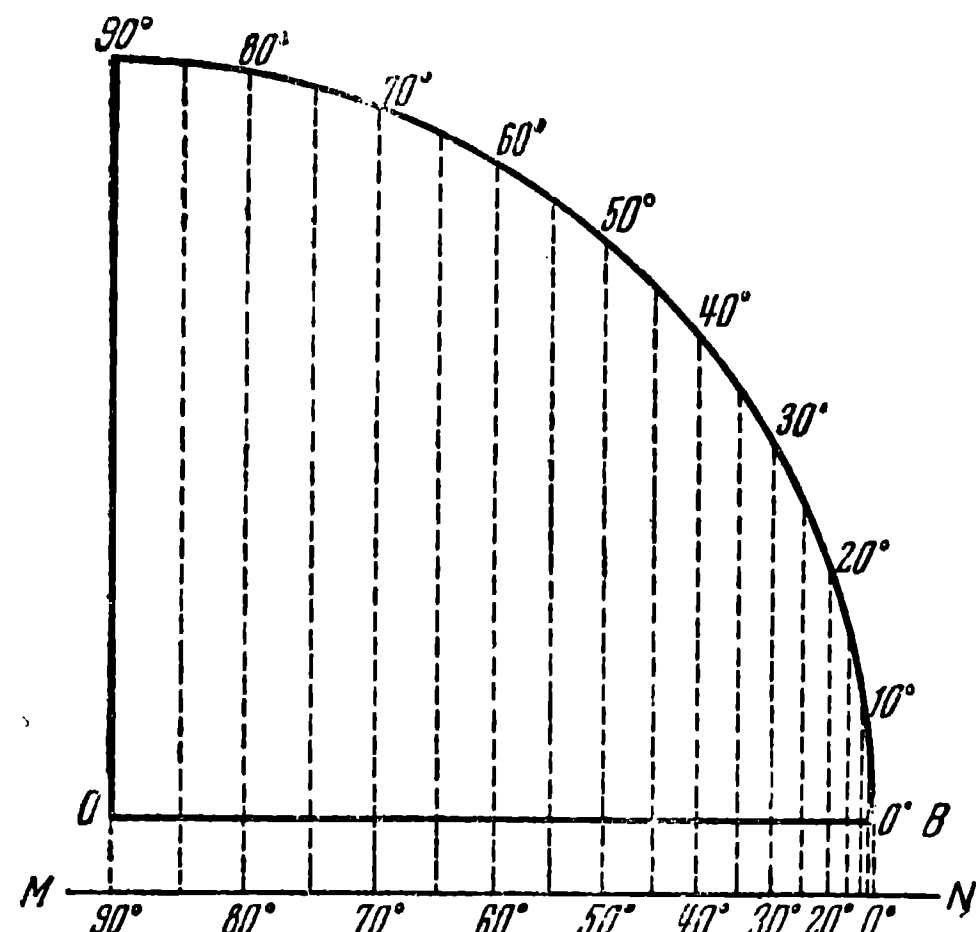


Рис. 26.

при определённых значениях другой величины (аргумента), мы получаем ряд соответственных значений обеих величин и изображаем эти результаты графически в прямоугольной системе координат. Так, например, зависимость между напряжением U на электродах и силой тока I в ионизованных газах выражается некоторой функцией

$$I = f(U),$$

которая в системе прямоугольных координат U (абсциссы) и I (ординаты) даёт характерную кривую тока насыщения (рис. 27); такие кривые получаются, как известно, в результате прямых измерений силы тока I в ионизованном газе при различных напряжениях U . На оси абсцисс отложим равномерную шкалу, полагая U равным 0, 50,

100, 150, ... единиц напряжения и, построив соответствующие ординаты, возьмём проекции их концов на ось I . В результате этого, обозначив точки деления на оси ординат соответствующими значениями I , мы получаем шкалу функции, определяющей зависимость I от U для ионизованного газа, хотя аналитическое выражение этой функции или значения её параметров нам могут оставаться неизвестными. Модуль этой функциональной шкалы определяется масштабом, взятым на оси ординат, и может быть вычислен на основании формулы (9).

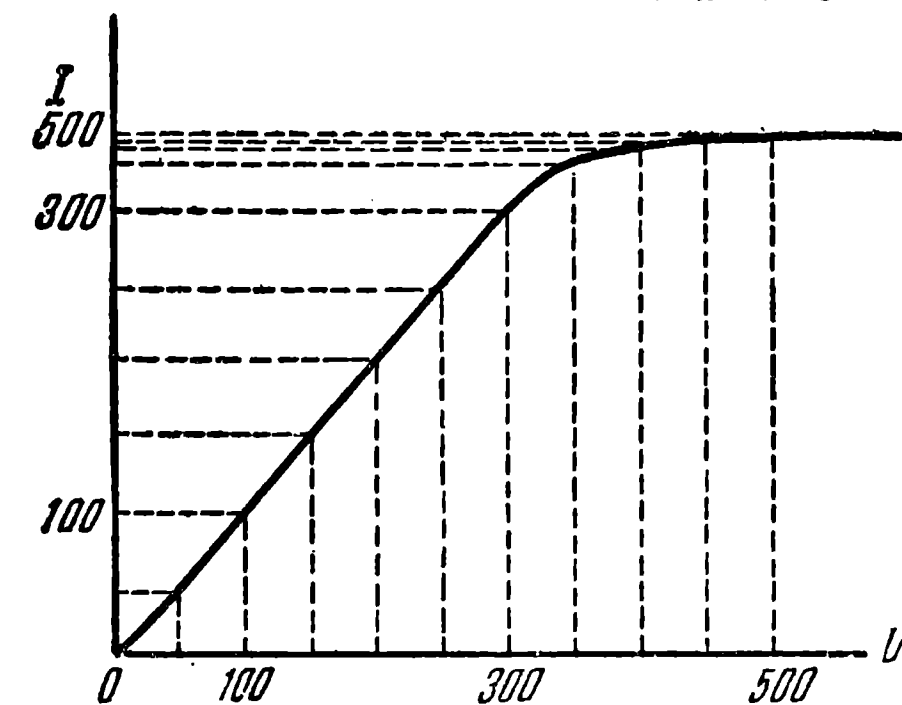


Рис. 27.

Иногда применяются функциональные шкалы с двумя системами делений, которые дают как значения функции y , так и значения аргумента x . Такие шкалы получили название **двойных шкал**. Двойную шкалу можно рассматривать как соединение шкалы равномерной и шкалы функциональной, причём их начальные и конечные точки должны совпадать. Двойная шкала косинуса в первом квадранте дана на рис. 28. Двойная шкала даёт возможность отсчитывать как значения функции по данным значениям аргумента, так и, наоборот, значения аргумента по данным значениям функции, т. е. может заменять таблицу значений функции, если не требуется большой точности.

Вместо двойной шкалы часто применяют систему двух шкал, функциональной и равномерной, которые строятся

на двух параллельных прямых так, чтобы начальные точки обеих шкал лежали на одной прямой, к ним перпендикулярной. Такая система двух шкал на параллельных прямых для значений косинуса в первом квадранте дана на рис. 29. При отсчёте в этом случае пользуются

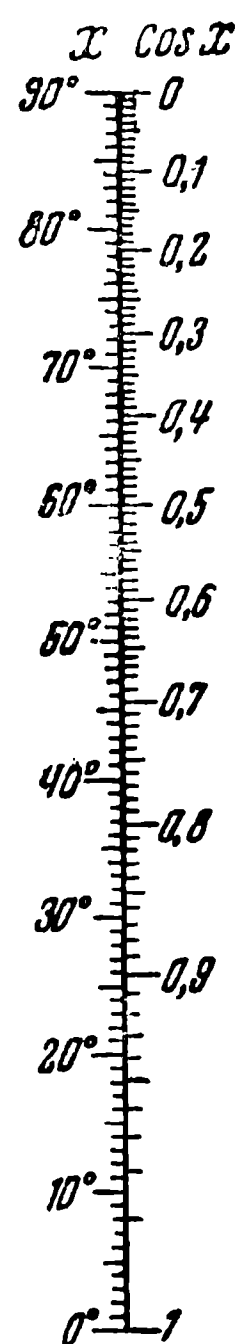


Рис. 28.

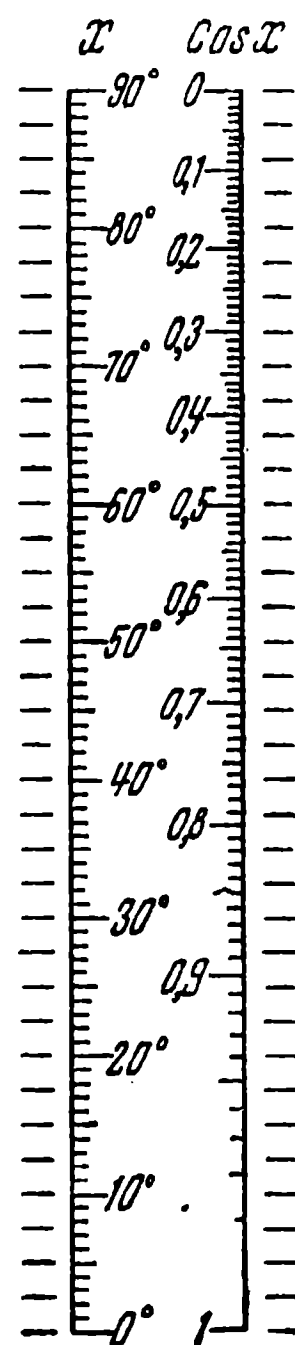


Рис. 29.

линейкой, предпочтительно прозрачной, которую направляют перпендикулярно к шкалам; для установления линейки иногда пользуются системой перпендикулярных к шкалам коротких прямых, которые намечаются по обе стороны шкал, как это показано на рисунке.

4. Логарифмическая шкала, степенная шкала и проективная шкала. Из различных функциональных шкал наибольшее применение при графических построении-

ях, кроме шкал тригонометрических функций, находят ещё л о г а р и ф м и ч е с к а я шкала, степенная шкала и проективная шкала. Основные свойства этих шкал можно определить следующим образом.

1) Л о г а р и ф м и ч е с к а я шкала. Логарифмическая шкала в простейшем случае определяется функцией

$$y = \lg x.$$

Для того чтобы построить логарифмическую шкалу, даём x ряд последовательных значений, равных, например, 1, 2, 3, ..., 10, и из таблиц берём логарифмы этих чисел, ограничиваясь вторым десятичным знаком. Значения x и y даны в таблице IX.

Таблица IX

x	1	2	3	4	5	6	7	8	9	10
y	0	0,30	0,48	0,60	0,70	0,78	0,84	0,90	0,95	1

Затем берём некоторый модуль, положив, например, $m=100$ мм, и откладываем значения y на прямой MN (рис. 30) от некоторой точки A , которую принимаем за

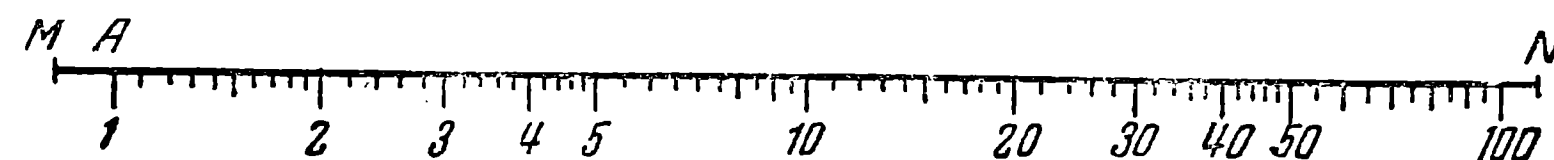


Рис. 30.

начало шкалы; против точек деления ставим соответствующие значения x . В результате мы получаем шкалу с неравномерными делениями, которая и представляет логарифмическую шкалу для интервала значений x от 1 до 10.

Возьмём ряд следующих значений x с интервалом, равным 10, т. е. положим x последовательно равным 10,

20, 30, ..., 100, и находим из таблиц соответствующие значения y . Значения x и y приведены в таблице X.

Таблица X

x	10	20	30	40	50	60	70	80	90	100
y	1,00	1,30	1,48	1,60	1,70	1,78	1,84	1,90	1,95	2,00

Откладывая эти значения y на логарифмической шкале, мы, очевидно, получим для интервала значений x от 10 до 100 расположение делений, в точности отвечающее тому, которое мы имеем в первом интервале, но эти деления мы должны нумеровать новыми значениями x , как это показано на рисунке. Точно так же, если мы возьмём ряд следующих значений x с интервалом, равным 100, т. е. положим x равным 100, 200, 300, ..., 1000, то вновь получим такое же расположение делений, которое обозначим новыми значениями x . Таким образом, мы видим, что логарифмическая шкала обладает особым свойством повторяемости, или воспроизведения делений, т. е. её деления через определённые интервалы в точности повторяются. Это свойство логарифмической шкалы является прямым следствием основного свойства логарифмов, согласно которому при изменении какого угодно числа в 10^n раз, где n —целое число, в десятичном логарифме этого числа изменяется только характеристика, а мантисса остаётся прежней. Это свойство логарифмической шкалы иногда коротко определяется так: деления логарифмической шкалы десятичных логарифмов в точности повторяются в интервалах между последовательными степенями десяти.

Логарифмическая шкала может быть построена и для более сложных выражений, например таких:

$$y = a + \lg x, \quad y = \lg(a + x), \quad y = \frac{1}{\lg(a + x)} \text{ и т. д.,}$$

где a обозначает некоторое постоянное число. Приёмы

построения шкалы в этих случаях совершенно аналогичны.

2) **Степенная шкала.** Степенная шкала в простейшем случае определяется функцией вида

$$y = x^m,$$

где m может быть целым или дробным числом; наиболее часто встречаются случаи, когда m равняется 2 и 3 или $\frac{1}{2}$ и $\frac{1}{3}$.

Степенная шкала, построение которой выполняется очень просто, обнаруживает резко выраженную нерав-

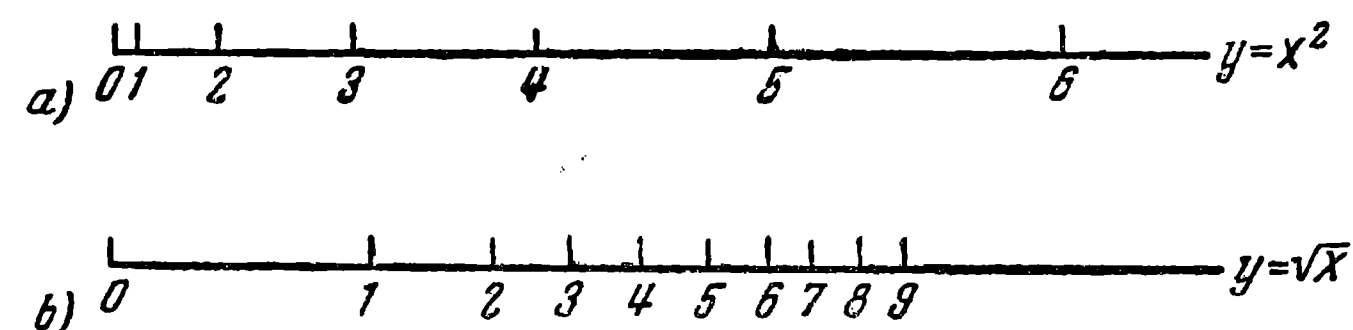


Рис. 31.

номерность делений, как это видно из рис. 31, где построены шкалы функций:

$$y = x^2 \text{ и } y = \sqrt{x}.$$

Эти шкалы построены для интервала значений x , первая от 0 до 6 и вторая от 0 до 9. Длина делений первой шкалы постепенно увеличивается, второй—уменьшается. Неравномерность степенной шкалы быстро возрастает при увеличении m . Вследствие большой неравномерности степенной шкалы она применяется обыкновенно только при относительно небольших интервалах изменения x .

3) **Проективная шкала.** Проективная шкала определяется функцией вида:

$$y = \frac{ax + b}{cx + d},$$

где a , b , c и d обозначают некоторые постоянные коэффициенты.

При построении проективной шкалы можно пользоваться обычным вычислительным приёмом, определяя значения y при данных значениях x . Но одновременно оказывается, что для построения проективной шкалы можно применять особый графический метод построения, на доказательстве которого мы останавливаться не будем. Этот метод заключается в следующем.

На прямой AB (рис. 32) берём равномерную шкалу и из некоторой точки O при помощи прямолинейных лучей проектируем деления этой шкалы на вторую прямую MN , наклонную по отношению к прямой AB . В результате

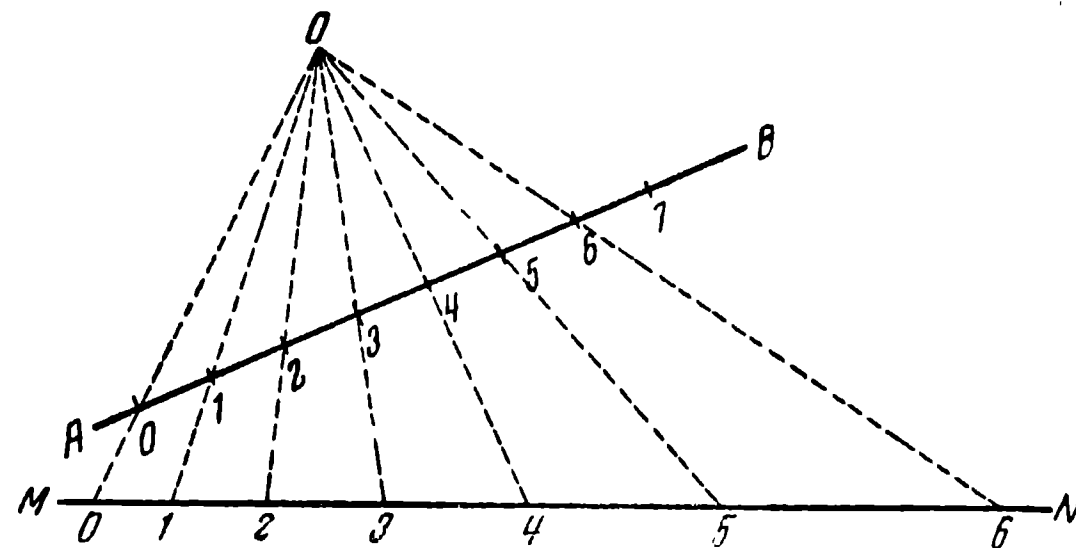


Рис. 32.

на прямой MN мы получаем шкалу с неравномерными делениями, которая представляет собой проективную шкалу. Относительное положение делений на ней зависит от положения центра O и угла между прямыми AB и MN .

Дальнейшим развитием проективной шкалы можно считать проективно-функциональную шкалу, которую мы получаем, если из некоторого центра проектируем на прямую MN не равномерную шкалу, а какую-либо функциональную шкалу, например степенную или тригонометрическую.

Все эти шкалы находят очень широкое применение при графических построениях, а также в различных приборах для вычислений; так, логарифмическая шкала, а также шкалы тригонометрических функций применяются, как известно, в логарифмической линейке.

5. Координатные сетки. При графических построениях очень часто применяются так называемые координатные сетки, которые в простейшем случае представляют собой системы равноотстоящих прямых, параллельных осям координат. Эти прямые мы проводим через точки деления на осях координат, на которых мы строим равномерные шкалы обыкновенно при одном масштабе на обеих осях. В результате этого вся плоскость чертежа оказывается разбитой на равновеликие квадраты, длина стороны которых определяется размерами делений на осях координат. Обыкновенная миллиметровая или координатная бумага может служить примером подобного рода координатных сеток, применение которых при графических построениях является чрезвычайно удобным.

Такую же сетку мы можем построить и в случае полярных координат; сетки для полярных координат также имеются в готовом виде.

Применяя для графиков готовую координатную бумагу, необходимо иметь в виду то, что по этому вопросу было сказано раньше (стр. 208).

Кроме равномерных координатных сеток, очень широкое применение находят функциональные координатные сетки, которые можно рассматривать как дальнейшее развитие равномерных сеток. Функциональные сетки мы получаем тем же приёмом, но на осях координат мы строим не равномерные, а функциональные шкалы; в общем случае они могут быть различны для оси x и для оси y . Если через точки деления этих шкал провести прямые, параллельные осям координат, то мы и получим функциональную сетку. Вся координатная плоскость оказывается при этом разделённой не на квадраты, а на прямоугольники, длина сторон которых в различных местах будет, очевидно, различной, так как размеры и форма прямоугольников определяются положением точек деления на осях координат, т. е. видом функциональных шкал, построенных на них.

Особенно широкое применение при графических построениях находят логарифмические функциональные сетки, при построении которых на обеих осях координат строятся логарифмические шкалы

(рис. 33). Логарифмические сетки приходится применять настолько часто, что изготавливается специальная л о г а р и ф м и ч е с к а я б у м а г а, на осях которой отложены логарифмические шкалы.

Часто применяется также так называемая п о л у л о г а р и ф м и ч е с к а я с е т к а; при её построении на одной из осей, например на оси x , строится логарифмическая шкала, а на другой оси строится равномерная шкала (рис. 34). Соответствующая полулогарифмическая бумага также имеется в готовом виде.

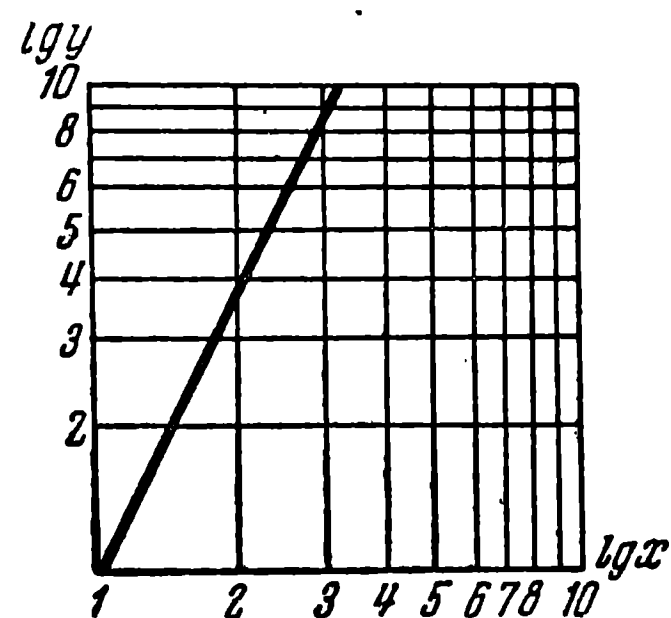


Рис. 33.

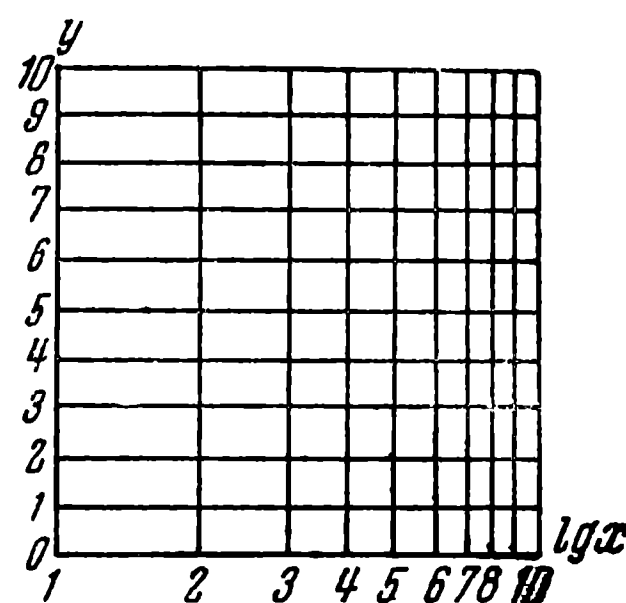


Рис. 34.

6. Спрямление кривых при помощи функциональных сеток. Если некоторую кривую

$$y = f(x)$$

мы будем изображать графически, применяя сначала равномерную сетку, а затем различные функциональные сетки, то, очевидно, форма кривой во всех этих случаях должна быть различной или, как обычно говорят, должно иметь место преобразование кривой в различные формы. Этим пользуются на практике, выбирая функциональные шкалы на осях координат такого вида, при котором данная кривая получает возможно более простую форму, например, преобразовывается в прямую линию.

Логарифмическая сетка даёт возможность ряд кривых преобразовывать в прямые линии или спрямлять

к р и в ы е, как часто говорят. Действительно, логарифмируя, например, уравнение

$$y = x^2,$$

находим:

$$\lg y = 2 \lg x.$$

Очевидно, что это уравнение в логарифмической сетке изображается прямой линией: значения её ординат для всех точек равны удвоенным значениям их абсцисс (рис. 33). К этим вопросам мы вернёмся ещё раз позднее (см. гл. X).

Г Л А В А VII

ЭЛЕМЕНТЫ НОМОГРАФИИ

1. Общее понятие о номографии. Иногда построение точных измерительных графиков в прямоугольных или полярных координатах оказывается очень сложным или не даёт удовлетворительных результатов. Это может объясняться, например, тем, что функциональная зависимость между переменными выражается очень сложной формулой, точное построение которой может быть весьма затруднительным. Или кривая может иметь такой вид, при котором графический отсчёт в некоторых интервалах графика не даёт достаточно точных результатов. Так, при отсчёте по гиперболе (рис. 17) трудно получить большую точность, определяя те значения y , которые отвечают небольшим значениям x , так как в этой области гипербола идёт почти параллельно оси ординат. Кроме того, как было сказано выше, графики, которые мы получаем обычными приёмами, можно применять для отсчётов только при тех частных значениях постоянных параметров, для которых графики были построены (стр. 211).

Вследствие всего этого постепенно были выработаны другие, более искусственные методы графического изображения функций. Эти методы дают возможность рассчитывать и оформлять такие графические построения, которые не утрачивают своего значения, т. е. могут служить для отсчётов, если параметры и коэффициенты в формулах изменяются, по крайней мере в некоторых пределах. Все эти вопросы в настоящее время разрабатываются в особом отделе прикладной математики, который получил название номографии. Этот термин в переводе с греческого

языка ($\nu\mu\omicron\varsigma$ —закон, $\gamma\rho\acute{\alpha}\phi\omega$ —чертить, рисовать) обозначает черчение или графическое изображение законов. Таким образом, в номографии разрабатываются методы графического воспроизведения различных математических выражений, например законов физических явлений, технических расчётных формул и т. д. В результате оказывается возможным вычертить или точно графически построить номограмму соответствующего выражения, что даёт возможность все вычислительные операции с данной формулой заменить чисто механическими отсчётами, обычно очень простыми. Эти отсчёты выполняются на глаз в большинстве случаев при помощи точной линейки или прямой линии, нанесённой на прозрачной целлулоидной пластинке.

Номограммой любой формулы можно пользоваться не при одном, а при различных значениях тех постоянных коэффициентов или параметров, которые входят в данную формулу, иначе говоря, номограмма любого выражения, однажды рассчитанная и точно вычерченная, сохраняет своё значение при изменении параметров в этом выражении. Таким образом, номограммы можно рассматривать до некоторой степени как изображение алгебраических соотношений между переменными величинами, которые остаются правильными при изменении числовых значений параметров в формулах, по крайней мере в некоторых пределах. Хотя расчёт и графическое оформление номограмм для некоторых формул часто оказываются очень кропотливой операцией, всё же эта затрата времени себя вполне оправдывает в тех случаях, когда необходимо вычислительные операции с этими формулами повторять очень много раз или обращаться к ним систематически.

Возможности и методы построения номограмм оказываются очень разнообразными, так что для одной и той же функциональной зависимости можно построить целый ряд различных номограмм. Вследствие этого введена определённая классификация номограмм, и в зависимости от метода построения их принято делить на отдельные виды и типы. Так, известны номограммы сетчатые, номограммы в выравненных точках, номограммы

с криволинейными шкалами и т. д. Из всего этого разнообразия номограмм мы ограничимся рассмотрением только наиболее простых видов номограмм, а именно, познакомимся с основными видами сетчатых номограмм и номограмм на выравненных точках.

2. Сетчатые номограммы. Различные типы номограмм и способы их построения удобнее всего разобрать на каком-либо простом примере. Возьмём уравнение изотерми-

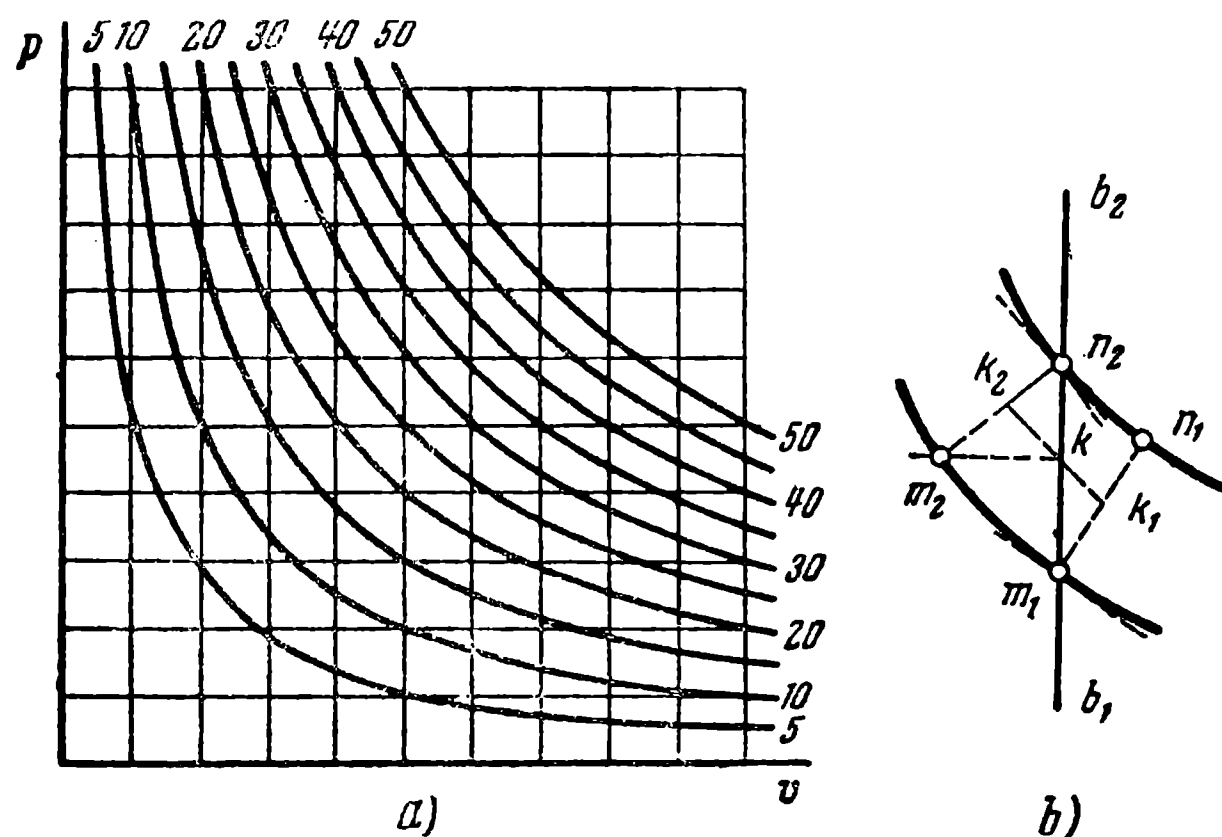


Рис. 35.

ческого процесса в газах или формулу умножения двух величин:

$$p \cdot v = a, \quad (1)$$

и попробуем построить для него четыре сетчатые номограммы наиболее простых типов. Для того чтобы можно было выполнить эти построения, допустим, что постоянный параметр a в этом уравнении изменяется в некоторых пределах, например от 1 до 100, а переменные величины p и v изменяются в одинаковых пределах, от 1 до 10. В этом предположении мы должны дать несколько номограмм уравнения (1), т. е. таких графических изображений этого уравнения, при помощи которых мы могли бы находить

значения p по данным значениям v или обратно, и притом при всех значениях a , которые лежат в указанных пределах. Одновременно эти номограммы должны давать возможность определить значения a в тех же пределах его изменения, если даны значения p и v .

Воспользуемся для этого следующими четырьмя приёмами.

1) На обеих осях берём одинаковые модули (стр. 221) и для значений a , взятых с интервалами, равными, например, 5 или 10, строим соответствующие гиперболы, как это изображено на рис. 35, а. График строят на равномерной координатной сетке: соответствующие значения a или индексы кривых ставят для некоторых гипербол обыкновенно с двух сторон вблизи их концов, которые несколько продолжают за рамку сетки. Такой график даёт возможность находить значения p по данным значениям v или обратно, причём постоянный параметр в уравнении (1) может иметь различные значения. Таким приёмом мы строим номограмму закона Бойля—Мариотта в виде семейства гипербол; при её применении на практике могут встретиться два случая.

а) Если данное нам значение параметра a отвечает одной из гипербол номограммы, то определение значений p по данным значениям v или обратно не представляет труда. При этом, если значение v не отвечает целым делениям шкалы на оси абсцисс или если получаются значения p , не отвечающие целому числу делений на оси ординат, то необходимо оценивать доли деления на глаз или применять графическое интерполирование (стр. 259). В данном случае оно выполняется очень просто, так как шкалы на обеих координатных осях являются равномерными, вследствие чего при определении p или v мы ограничиваемся простым отсчётом их значений по шкалам на координатных осях с точностью до 0,1 деления шкалы.

б) Если значение параметра a , равное a_k , не соответствует ни одной из гипербол номограммы, то при определении p или v также необходимо применять интерполирование, которое, однако, в данном случае оказывается более сложным и обычно выполняется следующим приёмом:

в точке абсциссы, отвечающей данному значению v , которое обозначим через V_k , восстанавливают ординату, продолжая её до пересечения с теми двумя гиперболами, между индексами которых лежит значение a_k . Такое построение сделано схематически на рис. 35, *b*, где m_1 и n_2 обозначают точки пересечения ординатой $b_1 b_2$ двух соседних гипербол. В этих точках проводим нормали к гиперболам и измеряем в произвольных единицах, например в миллиметрах, длину отрезков нормалей $m_1 n_1$ и $m_2 n_2$, т. е. расстояния между гиперболами. Допуская, что в пределах того небольшого интервала значений a , который отвечает двум соседним гиперболам, расстояния между гиперболами изменяются пропорционально их индексам, находим при помощи линейного интерполирования те точки k_1 и k_2 , которые отвечают гиперболе с индексом a_k , иначе говоря, мы допускаем, что если вычертить гиперболу с индексом a_k , то она должна пройти через точки k_1 и k_2 . При таком определении точек k_1 и k_2 мы, конечно, допускаем некоторую ошибку, которая, однако, будет невелика, если гиперболы лежат близко одна к другой, т. е. если интервалы значений a взяты достаточно малыми. Точки k_1 и k_2 соединяем вместо отрезка гиперболы отрезком прямой линии, которая пересекает ординату $b_1 b_2$ в точке k . Это построение вновь вносит некоторую ошибку, также небольшую при малых расстояниях между гиперболами. Проектируя точку k на ось ординат, отсчитываем на ней то давление p_k , которое отвечает данному объёму v_k .

Если значение a лежит за пределами взятого для него интервала, то номограмма не утрачивает своего значения и ею попрежнему можно пользоваться для определения значений p или v .

В этом случае следует изменить соответствующим образом цену делений на осях координат, т. е. пределы изменения p и v , так, чтобы значение a вновь оказалось внутри того интервала его значений, который определяется крайними значениями p и v .

Нетрудно видеть, что эта номограмма даёт возможность находить значения a по данным значениям p и v , т. е. является номограммой умножения чисел. Для этого

следует построить точки, отвечающие данным значениям p и v , и определить затем индексы соответствующих гипербол, применяя интерполирование, если это оказывается необходимым.

Номограммы подобного типа можно построить для любого уравнения с тремя неизвестными, т. е. для функции

$$f(x, y, z) = 0, \quad (2)$$

если рассматривать два из неизвестных как координаты точек на плоскости, а третье неизвестное считать параметром. Изменяя значение параметра в некоторых пределах, мы получим номограмму в виде семейства кривых линий, форма которых определяется видом функции f . Такие номограммы принято называть сетчатыми номограммами. Применение криволинейных сетчатых номограмм на практике неизбежно ограничивается чрезвычайной сложностью их построения. Поэтому сетчатые криволинейные номограммы заменяют, если это возможно, номограммами прямолинейными, пользуясь различными приёмами спрямления кривых.

2) Возможность заменить криволинейную сетчатую номограмму прямолинейной иногда определяется видом функции f в уравнении (2), что имеет место и в случае закона Бойля—Мариотта.

Напишем уравнение (1) в таком виде:

$$v = \frac{1}{p} \cdot a. \quad (3)$$

В этом уравнении будем считать переменными величинами v и a , а параметром p . Его значения будем по-прежнему изменять в пределах от 1 до 10, причём для удобства дальнейших построений вычислим величины, обратные p ; результаты этих вычислений приведены в таблице XI.

Таблица XI

p	1	2	3	4	5	6	7	8	9	10
$\frac{1}{p}$	1,00	0,50	0,33	0,25	0,20	0,17	0,14	0,12	0,11	0,10

В соответствии с этим из уравнения (3) получаем:

$$\begin{aligned} v &= a & v &= 0,17a, \\ v &= 0,50a, & v &= 0,14a, \\ v &= 0,33a, & v &= 0,125a, \\ v &= 0,25a, & v &= 0,11a, \\ v &= 0,20a, & v &= 0,1a. \end{aligned}$$

Если эти уравнения изобразить графически в системе прямоугольных координат, откладывая значения v по оси абсцисс, а значения a по оси ординат, то получается

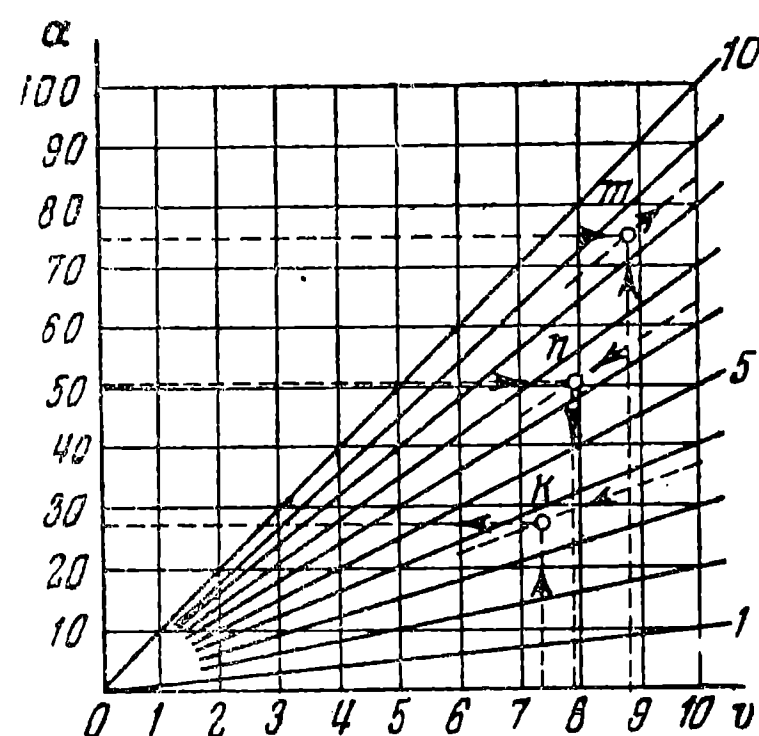


Рис. 36.

семейство прямых линий, или лучей, которые проходят через начало координат, причём значения $1/r$ служат их угловыми коэффициентами (рис. 36). При построении этой номограммы мы пользуемся равномерными шкалами на обеих осях, но модуль шкалы на оси абсцисс берём в десять раз больше модуля шкалы на оси ординат; отношение модулей на координатных осях мы выбираем в зависимости от относительных значений a и v , которые в данном случае относятся как 10 к 1. Лучи строятся на равномерной прямоугольной сетке, их концы, которые сходятся в начале координат, обыкновенно несколько не доводят до него, чтобы не иметь пересечения большого числа прямых в одной точке. Вторые концы лучей нумеруют значениями r , как это показано на рисунке.

В результате такого построения мы получили новый вид номограммы закона Бойля—Мариотта, или номограммы умножения чисел. Номограммы такого вида относятся также к числу сетчатых номограмм, но прямолинейных;

эти номограммы в соответствии с их видом принято называть иногда *лучевыми*, или *радиантными* номограммами. Очевидно, что построение лучевых номограмм выполняется очень просто по сравнению с построением криволинейных номограмм, так как все лучи проходят через начало координат, и для того чтобы их построить, достаточно определить ещё только одну точку для каждого луча.

Применение этой номограммы для определения значений одной из переменных величин в уравнении (1) по данным значениям двух других очень несложно.

Действительно, различные случаи, которые здесь возможны, ограничиваются, очевидно, следующими:

а) При определённом значении величины a , которую мы при отсчётах по номограмме продолжаем считать параметром, дано значение v . Для того чтобы определить значение r , отвечающее этим данным, строим точку с координатами a и v . Если она не совпадает ни с одним из лучей номограммы, то намечаем луч, отвечающий этой точке, и отсчитываем ординату точки его пересечения с крайней правой ординатой графика. Такое построение показано на номограмме (точка m) в предположении, что a равно 75 и v равно 8,8.

б) При определённом значении a дано значение r . Для того чтобы определить соответствующее значение v , откладываем на крайней правой ординате номограммы значение r , а на оси ординат—значение a . Если их крайние точки не соответствуют ни одному из делений, то намечаем соответствующий луч и прямую, параллельную оси v ; абсцисса точки их пересечения отвечает искомому значению v . Такое построение показано на номограмме (точка n) в предположении, что a равно 50,1 и r равно 6,4.

в) Наконец, если даны значения v и r , то для определения a намечаем абсциссу и луч, отвечающие данным значениям v и r ; ордината точки пересечения этих прямых даёт искомое значение a . Такое построение показано на номограмме (точка k) в предположении, что r равно 3,7 и v равно 7,3.

3) Лучевую номограмму, подобную описанной, можно получить в данном случае, применяя ещё другой метод

спрямления гиперболы. Для этого уравнение (3) напишем в таком виде:

$$\frac{1}{p} = \frac{1}{a} \cdot v,$$

и введём новую переменную, полагая

$$z = \frac{1}{p}. \quad (4)$$

Имеем:

$$z = \frac{1}{a} \cdot v.$$

Если это уравнение изобразить графически в прямоугольной системе координат, откладывая по оси абсцисс значения v , а по оси ординат значения z , то мы получим семейство прямых линий или лучей, проходящих через начало координат, для которых значения $1/a$ служат угловыми коэффициентами.

Для того чтобы построить соответствующую номограмму, необходимо определить вид шкалы на обеих осях координат и найти значения углового коэффициента лучей.

На оси абсцисс, мы, очевидно, должны оста-

вить равномерную шкалу, которую строим для значений v в прежнем масштабе (рис. 37). На оси ординат следует отложить значения z , т. е. значения величины, обратной p ; таким образом, мы видим, что на оси ординат следует построить шкалу функции, определяемой уравнением (4). Для этого даём величине p ряд значений в границах интервала его изменений, т. е. полагаем, что p равно 1, 2, 3, ..., 10, и вычисляем соответствующие значения z . Получаем результаты, приведённые в таблице XII.

Эту шкалу строим на оси ординат, отмечая её деления соответствующими значениями p : при построении этой шкалы мы выбираем такой модуль, чтобы общая

Таблица XII

p	1	2	3	4	5	6	7	8	9	10
z	1,00	0,50	0,33	0,25	0,20	0,17	0,14	0,125	0,111	0,10

длина шкалы была равна той, которой мы пользовались раньше. Наличие на оси ординат функциональной шкалы с неравномерными делениями вызывает необходимость вычислить и нанести на ней ряд промежуточных делений, отвечающих дробным значениям p , а затем построить функциональную сетку, ограничиваясь в данном случае системой прямых, параллельных оси абсцисс, как это показано на рисунке. После этого находим угловые коэффициенты всех прямых и строим лучи, отмечая их соответствующими значениями a . В результате мы получаем сетчатую прямолинейную номограмму формулы (1) несколько иного вида, применение которой для определения значений p , v и a нетрудно разобрать по аналогии с предыдущей номограммой. Вычисления значений p , v и a в данном случае оказываются несколько сложнее вследствие неравномерной шкалы на оси ординат, что осложняет интерполирование и вызывает необходимость применять более частую сетку в горизонтальном направлении.

4) Наконец, прямолинейную номограмму уравнения (3) можно ещё получить, переходя к логарифмическим шкалам. Для этого логарифмируем уравнение (1):

$$\lg v + \lg p = \lg a.$$

Введём новые переменные, полагая

$$x = \lg v \quad \text{и} \quad y = \lg p. \quad (5)$$

Находим:

$$x + y = \lg a.$$

Если это уравнение изобразить графически, откладывая по оси абсцисс значения x , а по оси ординат значения y , то получится семейство прямых, которые отсекают на осях координат равные отрезки, т. е. семейство прямых, которые образуют с осями координат углы, равные 45° . Для того чтобы построить эту номограмму, строим на осях координат в соответствии с уравнением (5) логарифмические шкалы с одним и тем же модулем, величину которого выбираем так, чтобы при взятом выше интервале значений p и v общая длина шкалы равнялась длине шкал в прежних номограммах. При построении шкал следует построить ряд промежуточных делений, которые отвечают дробным значениям p и v , деления шкал помечаем соответствующими значениями p и v , как это показано на рис. 38.

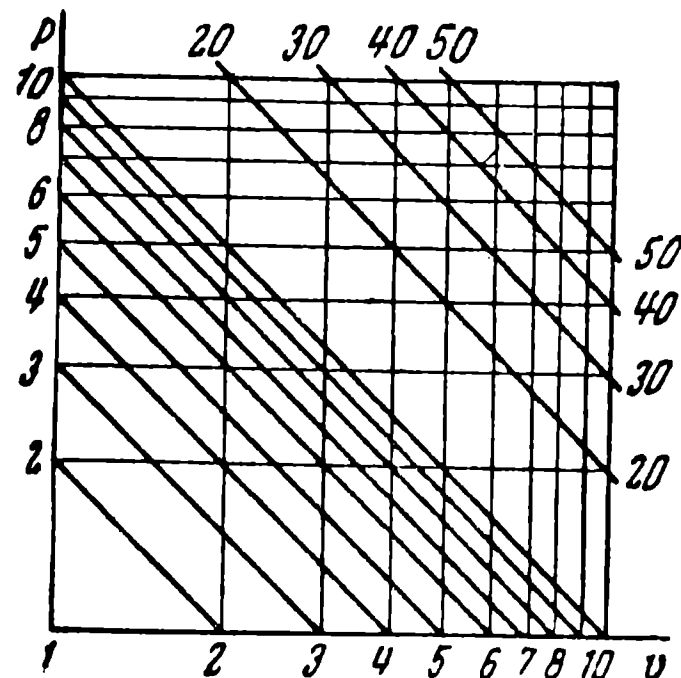


Рис. 38.

Затем строим логарифмическую сетку и, соединяя прямыми линиями деления обеих шкал, имеющие одинаковые индексы, получаем новый вид сетчатой прямолинейной номограммы формулы (1). В её применении для определения значений p , v и a нетрудно разобраться. Так, например, если при определённом значении a дано значение v , то для определения p следует отложить на оси абсцисс данное значение v и, восстановив в этой точке перпендикуляр, определить ординату точки его пересечения с той прямой, которая отвечает значению a .

Такое преобразование криволинейных номограмм в прямолинейные, основанное на применении логарифмических шкал, обычно называется логарифмической анаморфозой номограммы.

Сетчатые номограммы прямолинейного вида находят достаточно широкое применение, но преобразование семейства кривых в прямые линии далеко не всегда является столь же простым, как в рассмотренном случае, и

анаморфоза номограмм, при которой криволинейные номограммы заменяются прямолинейными, оказывается возможной далеко не всегда.

3. Номограммы на выравненных точках. Вследствие большой простоты и удобства применения номограммы на выравненных точках являются чрезвычайно распространёнными. Эти номограммы в настоящее время представлены очень большим числом различных типов; из них

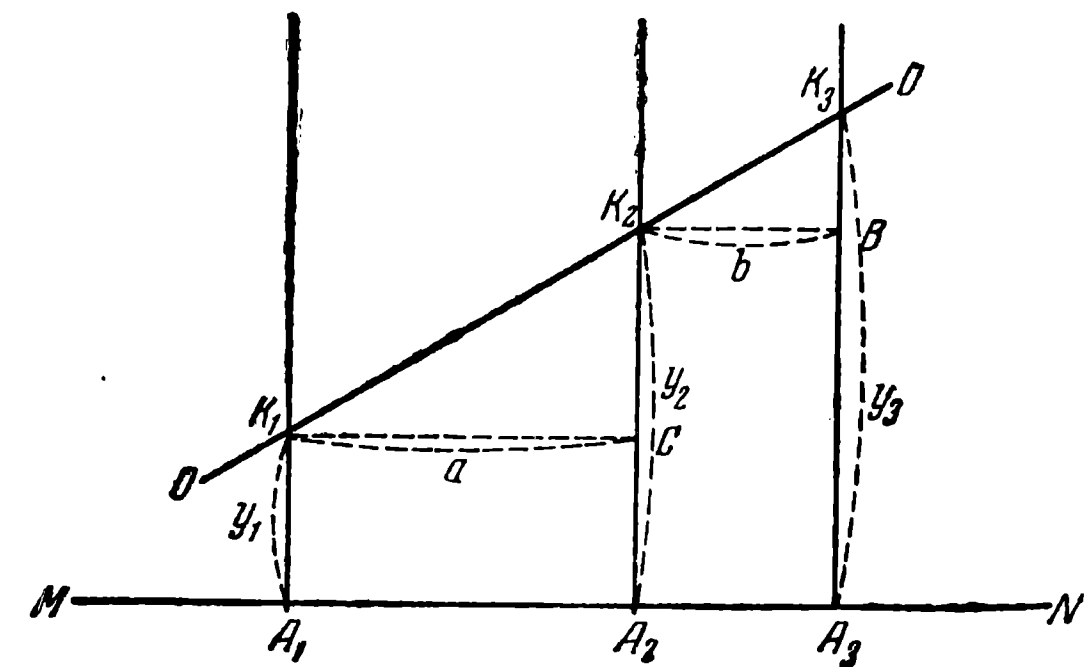


Рис. 39.

наиболее простыми можно считать номограммы с тремя параллельными шкалами и зет-номограммы, рассмотрением которых мы и ограничимся.

В данном параграфе мы рассмотрим номограммы с тремя параллельными шкалами. Возьмём три функциональные шкалы, отвечающие уравнениям

$$\begin{aligned} y_{1i} &= f_1(x_1), \\ y_{2i} &= f_2(x_2), \\ y_{3i} &= f_3(x_3). \end{aligned} \quad (6)$$

Мы предполагаем, что эти шкалы построены на трёх параллельных прямых (рис. 39) и все три направлены в одну сторону, причём начальные точки шкал A_1 , A_2 и A_3 выбираем на прямой MN , перпендикулярной к направлению шкал. Обозначая модули шкал соответственно

через m_1 , m_2 и m_3 , мы можем написать (стр. 220):

$$\begin{aligned}\bar{y}_{1i} &= m_1 f_1(x_1), \\ \bar{y}_{2i} &= m_2 f_2(x_2), \\ \bar{y}_{3i} &= m_3 f_3(x_3).\end{aligned}\quad (6')$$

Такое построение может служить номограммой некоторой функциональной зависимости между переменными x_1 , x_2 и x_3 , если мы определим необходимые для этого условия. С этой целью возьмём произвольную прямую OO , которая пересекает шкалы в точках K_1 , K_2 и K_3 ; расстояния этих точек от начальных точек шкал обозначим соответственно через $\bar{y}_1$, $\bar{y}_2$ и $\bar{y}_3$. Из подобия треугольников K_1K_2C и K_2K_3B имеем:

$$\frac{\bar{y}_2 - \bar{y}_1}{\bar{y}_3 - \bar{y}_2} = \frac{K_1C}{K_2B} = \frac{a}{b}.$$

Отсюда находим:

$$\bar{y}_2 = \frac{b\bar{y}_1 + a\bar{y}_3}{a + b}.$$

Из этого выражения на основании равенств (6) и (6') получаем:

$$f_2(x_2) = \frac{m_1 b f_1(x_1) + m_3 a f_3(x_3)}{m_2(a + b)}.$$

Обозначая коэффициенты при $f_1(x_1)$ и $f_3(x_3)$ в правой части этого уравнения через q_1 и q_2 , имеем:

$$f_2(x_2) = q_1 f_1(x_1) + q_2 f_3(x_3), \quad (7)$$

где q_1 и q_2 определяются выражениями

$$q_1 = \frac{m_1 b}{m_2(a + b)} \quad \text{и} \quad q_2 = \frac{m_3 a}{m_2(a + b)}. \quad (8)$$

Выражение (7) определяет зависимость между переменными x_1 , x_2 и x_3 , которая устанавливается номограммой данного вида. Эта номограмма, как показывает вывод уравнения (7), даёт возможность очень просто определять значения одной из трёх переменных x_1 , x_2 и x_3 , если значения двух других даны. Действительно, если, например, даны значения x_1 и x_3 и требуется определить зна-

чение x_2 , то достаточно соединить прямой линией те деления на первой и третьей шкалах, отметки которых отвечают данным значениям x_1 и x_3 ; точка пересечения этой прямой со средней шкалой даёт нам искомое значение x_2 . Совершенно так же можно определять значения x_1 или x_3 , если даны значения x_2 и x_3 или x_1 и x_2 . Отсюда становится понятным название этого типа номограмм — номограммы на выравненных точках.

Вполне справедливым оказывается и обратное заключение, а именно, мы вправе утверждать, что если между тремя переменными величинами x_1 , x_2 и x_3 имеет место функциональная зависимость, которая может быть представлена выражением вида (7), то для этих трёх величин можно построить номограмму данного типа, т. е. три функциональные шкалы на трёх параллельных прямых. При этом вид всех трёх функций $f_1(x_1)$, $f_2(x_2)$ и $f_3(x_3)$ и пределы изменения их аргументов должны быть даны; кроме того, для построения номограммы, как показывают уравнения (6) и (8), должны быть известны, во-первых, модули m_1 , m_2 и m_3 всех трёх шкал и, во-вторых, расстояния a и b между шкалами. Определение всех этих величин принято делать в определённой последовательности, а именно:

1) Определяют модули m_1 и m_3 тех двух шкал, которые отвечают функциям $f_1(x_1)$ и $f_3(x_3)$ в уравнениях (6). Вычисление этих модулей производится на основании уравнений (6'), принимая во внимание пределы изменения аргументов x_1 и x_3 и выбрав предварительно определённую длину шкал, обыкновенно одинаковую для той и другой шкал.

2) Затем определяют модуль третьей шкалы, что можно сделать, пользуясь уравнениями (8). Из этих уравнений находим:

$$\frac{q_1 m_2}{m_1} = \frac{b}{a + b} \quad \text{и} \quad \frac{q_2 m_2}{m_3} = \frac{a}{a + b}.$$

Складывая эти уравнения, после небольших преобразований получаем:

$$\frac{m_2(q_1 m_3 + q_2 m_1)}{m_1 m_3} = 1$$

или

$$m_2 = \frac{m_1 m_3}{q_1 m_3 + q_2 m_1}, \quad (9)$$

т. е. уравнение, из которого можно определить m_2 , если известны q_1 и q_2 , т. е. коэффициенты в уравнении (7). Для того чтобы найти эти величины, следует данное нам уравнение привести к виду (7) и в полученном уравнении значения коэффициентов при $f_1(x_1)$ и $f_3(x_3)$ принять равными соответственно величинам q_1 и q_2 .

3) Сумма величин $a + b$, т. е. расстояние между крайними шкалами, является произвольной, и мы выбираем её значения, исходя из общих условий при построении данной номограммы. Но произвольной является только сумма величин a и b , что же касается положения средней шкалы, то оно подчиняется определённым условиям, которое вытекает из уравнений (8). Действительно, разделив второе из этих уравнений на первое, находим:

$$\frac{a}{b} = \frac{q_2 m_1}{q_1 m_3}. \quad (10)$$

Таким образом, одну из величин, a или b , можно выбирать произвольно, значение второй величины определяется из уравнения (10).

Номограммы этого типа принято называть номограммами на выравненных точках с тремя параллельными шкалами.

Для того чтобы лучше освоить построение таких номограмм, возьмём попрежнему закон Бойля—Мариотта, т. е. уравнение (1), и построим для него номограмму подобного типа.

Уравнение (7) показывает, что такие номограммы можно применять в тех случаях, когда мы имеем сумму функций. Поэтому уравнение (3) мы логарифмируем:

$$\lg p + \lg v = \lg a. \quad (11)$$

Сравнивая это уравнение с уравнением (7), мы видим, что

$$\begin{aligned} \lg p &= f_1(x_1), \\ \lg v &= f_3(x_3), \\ \lg a &= f_2(x_2). \end{aligned}$$

Таким образом, на всех трёх прямых следует построить логарифмические шкалы. Будем считать, что даны значения p и v , а значения a подлежат определению. Поэтому удобнее на крайних прямых строить шкалы p и v , пользуясь прежними интервалами их значений, т. е. от 1 до 10. Для этих шкал берём одинаковый модуль: его значение выбираем так, чтобы получить ту же длину шкал для p и v , как и в прежних номограммах. Считая длину этих шкал равной, например, L мм, находим:

$$m_1 = m_3 = \frac{L \text{ мм}}{\lg 10 - \lg 1} = L \text{ мм}.$$

Такие шкалы строим на крайних прямых (рис. 40). Среднюю шкалу в данном случае следует расположить на прямой, которая находится на середине расстояния между крайними шкалами, что является непосредственным следствием полной симметричности уравнения (11) относительно p и v . Это следует также из того, что при непосредственном сопоставлении уравнений (11) и (7) мы находим, что

$$q_1 = q_2 = 1.$$

Таким образом, принимая $a + b$ равным, например, 60 мм, имеем:

$$a = b = 30 \text{ мм}.$$

Поэтому, вычисляя модуль средней шкалы по формуле (9), находим:

$$m_2 = \frac{m_1 m_3}{m_1 + m_3} = \frac{L}{2},$$

т. е. при построении средней шкалы для $\lg a$ следует взять модуль, в два раза меньший по сравнению с модулем крайних шкал. Этой номограммой мы можем

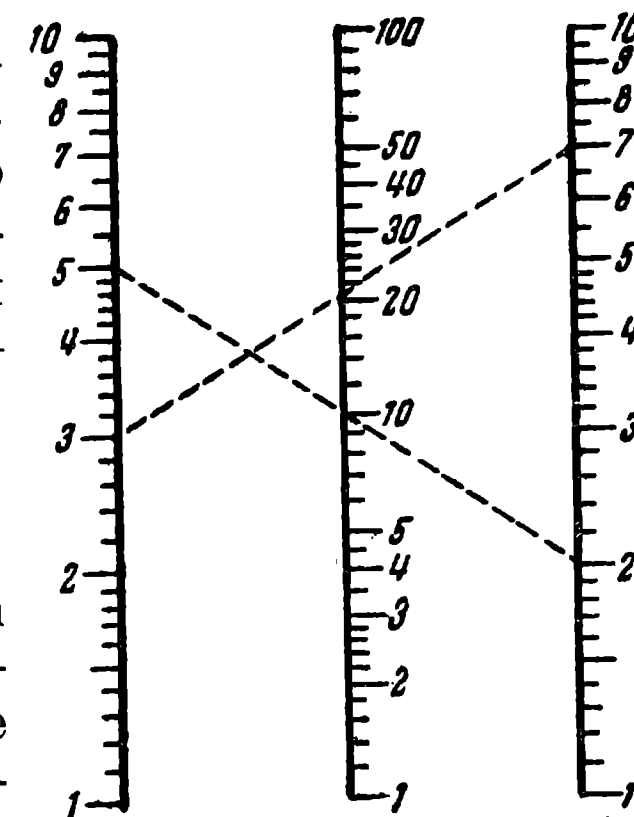


Рис. 40.

пользоваться, очевидно, для определения какой угодно из трёх переменных величин, входящих в уравнение (3), если даны значения двух других. Способ пользования номограммой понятен из предыдущего.

4. Зет-номограммы. Возьмём две функциональные шкалы, отвечающие уравнениям

$$y' = f_1(x_1),$$

$$y'' = f_2(x_2).$$

Эти шкалы мы строим на двух параллельных прямых A_1B_1 и A_2B_2 (рис. 41), причём мы предполагаем, что эти

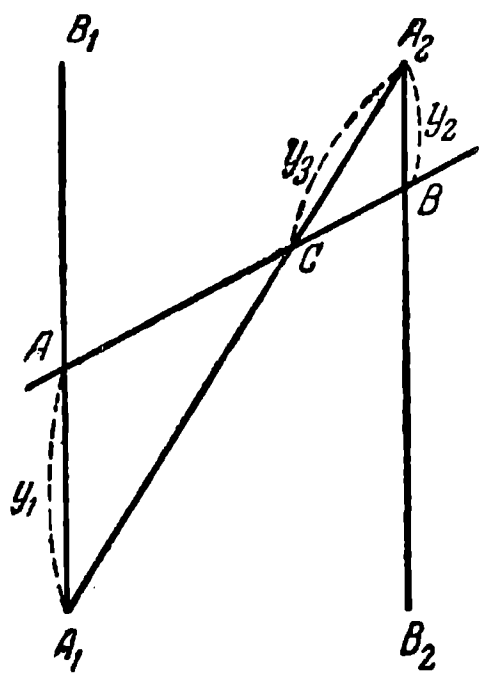


Рис. 41.

шкалы направлены в противоположные стороны, так что начальные точки шкал находятся в точках A_1 и A_2 . Третью функциональную шкалу, отвечающую уравнению

$$y''' = f_3(x_3),$$

построим на прямой A_1A_2 , которая соединяет начальные точки первой и второй шкал, причём начало третьей шкалы мы помещаем в точке A_2 . Определим характер той зависимости между переменными x_1 , x_2 и x_3 , при которой такое построение может служить номограммой. Для этого пере-

сечём все три шкалы некоторой прямой AB и введём обозначения:

$$A_1A = \bar{y}_1, \quad A_2B = \bar{y}_2, \quad A_2C = \bar{y}_3 \quad \text{и} \quad A_1A_2 = h.$$

Из подобия треугольников A_1AC и A_2BC находим:

$$\frac{\bar{y}_1}{\bar{y}_2} = \frac{h - \bar{y}_3}{\bar{y}_3}. \quad (12)$$

Обозначая модули первой, второй и третьей шкал соответственно через m_1 , m_2 и m_3 , имеем:

$$\bar{y}_1 = m_1 f(x_1),$$

$$\bar{y}_2 = m_2 f(x_2),$$

$$\bar{y}_3 = m_3 f(x_3).$$

Вводя эти обозначения в уравнение (12), получаем:

$$\frac{f_1(x_1)}{f_2(x_2)} = \frac{m_2(h - m_3 f(x_3))}{m_1 m_3 f(x_3)}.$$

Правая часть этого уравнения зависит только от третьего переменного x_3 . Поэтому можно положить:

$$\frac{f_1(x_1)}{f_2(x_2)} = \varphi(x_3)$$

или

$$f_1(x_1) = f_2(x_2) \cdot \varphi(x_3), \quad (13)$$

где

$$\varphi(x_3) = \frac{m_2(h - m_3 f_3(x_3))}{m_1 m_3 f_3(x_3)}. \quad (14)$$

Решая это уравнение относительно $f_3(x_3)$, находим:

$$f_3(x_3) = \frac{h}{m_3 \left(1 + \frac{m_1}{m_2} \varphi(x_3) \right)}.$$

Вставляя это значение $f_3(x_3)$ в третье из уравнений (12), получаем:

$$\bar{y}_3 = \frac{h}{1 + \frac{m_1}{m_2} \varphi(x_3)}, \quad (15)$$

т. е. выражение, которое даёт возможность построить шкалу третьей функции $f_3(x_3)$. Ввиду сходства общего очертания этих номограмм они с буквой Z и получили название зет-номограмм.

Уравнение (13) определяет ту функциональную зависимость между переменными x_1 , x_2 и x_3 , при которой выполненное нами построение может служить номограммой. Мы вправе сделать и обратное заключение, т. е. если между тремя переменными величинами x_1 , x_2 и x_3 имеется зависимость, которая может быть представлена уравнением вида (13), то для этих величин можно построить зет-номограмму.

Но при этом необходимо иметь в виду, что и построение зет-номограмм осложняется в тех случаях, когда начальные точки шкал лежат далеко за пределами

рабочей части номограммы, так как в этих случаях соединить прямой линией начала шкал иногда оказывается очень трудным.

Разобранными видами номограмм далеко не исчерпываются возможности, которыми пользуется номография в настоящее время. Так, наряду с равномерными и функциональными шкалами, построенными на прямых линиях, в номографии очень широко применяются такие же шкалы, но построенные на кривых, или криволинейные шкалы, которые для некоторых функциональных зависимостей оказываются безусловно необходимыми; этих вопросов мы здесь касаться не будем, ограничиваясь теми элементарными сведениями, которые были изложены выше.

Номограммы обыкновенно вычерчиваются в достаточно большом масштабе, длину шкал обычно берут в среднем около 25 или 30 см, иногда ещё больше. При таких условиях отсчёт по номограмме даёт достаточно точные результаты; их точность в среднем соответствует той, которая получается при вычислениях с логарифмической линейкой или с трёхзначными таблицами логарифмов.

5. Примеры расчёта и построения номограмм. Рассмотрим несколько примеров расчёта и построения номограмм.

Пример 1. Момент инерции J бесконечно тонкого прямоугольника (рис. 42) с основанием, равным d , и высотой, равной h , относительно оси O_1O_2 , проходящей через его центр O параллельно высоте, определяется выражением

$$J = \frac{h d^3}{12}. \quad (16)$$

Построим для этого выражения номограмму на выравненных точках на трёх параллельных шкалах, предполагая, что d и h изменяются в одинаковых пределах от 1 до 10.

Логарифмируя формулу (16), находим:

$$\lg J + \lg 12 = \lg h + 3 \lg d.$$

Сравнивая это уравнение с уравнением (7), мы видим, что

$$\lg J + \lg 12 = f_2(x_2),$$

$$\lg h = f_1(x_1),$$

$$\lg d = f_3(x_3)$$

и, кроме того,

$$q_1 = 1 \text{ и } q_2 = 3.$$

Таким образом, на всех прямых следует построить логарифмические шкалы. На крайних прямых построим шкалы $\lg h$ и $\lg d$, причём для них возьмём равные модули, т. е. положим:

$$m_1 = m_3 = m.$$

Если общую длину этих шкал мы примем равной 20 см, то, вычисляя их модули, находим:

$$m = \frac{200 \text{ мм}}{\lg 10 - \lg 1} = 200 \text{ мм}.$$

На основании уравнения (10) найдём отношение a и b :

$$\frac{a}{b} = \frac{m_1 q_2}{m_3 q_1} = \frac{q_2}{q_1} = 3 : 1,$$

т. е. среднюю шкалу следует поместить на одной четверти расстояния между шкалами, ближе к шкале $\lg d$. Модуль средней шкалы находим по формуле (9):

$$m_2 = \frac{m_1 m_3}{1 \cdot m_3 + 3 m_1} = \frac{m}{4} = 50,$$

т. е. модуль третьей шкалы следует взять равным 50 мм. Начальную точку для этой шкалы выбираем так, чтобы три точки, удовлетворяющие формуле (16), располагались на одной прямой, например точки, отвечающие значениям

$$J = 2, \quad h = 3 \text{ и } d = 2.$$

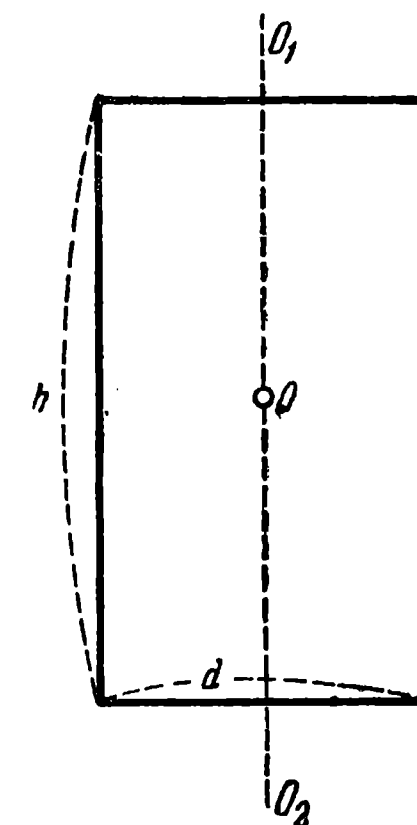


Рис. 42.

Номограмма дана на рис. 43 в масштабе 1:2 по отношению к натуральной величине; расстояние между крайними шкалами взято в два раза меньше их длины.

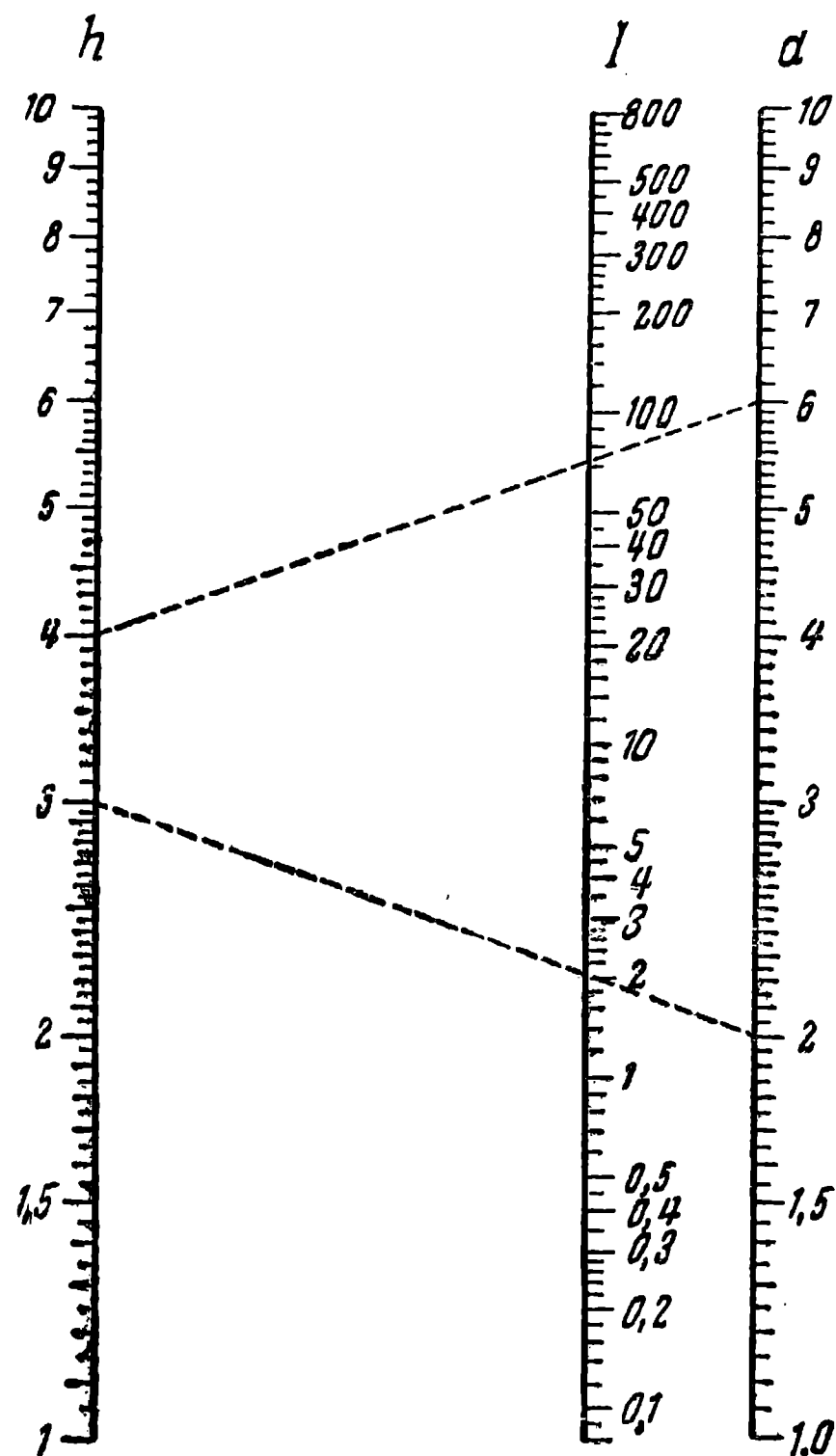


Рис. 43.

Пример 2. Построим вет-номограмму степенной функции:

$$y = x^n, \quad (17)$$

предполагая, что переменные изменяются в пределах

y от 1 до 10, x от 1 до 100 и n от 0 до 10.

Логарифмируя выражение (17), находим:

$$\lg y = n \lg x.$$

Сравнивая это уравнение с уравнением (13), мы видим, что

$$\lg y = f_1(x_1), \quad \lg x = f_2(x_2), \quad n = \varphi(x_3).$$

Шкалы функций $\lg y$ и $\lg x$ строим на двух параллельных прямых в противоположных направлениях,

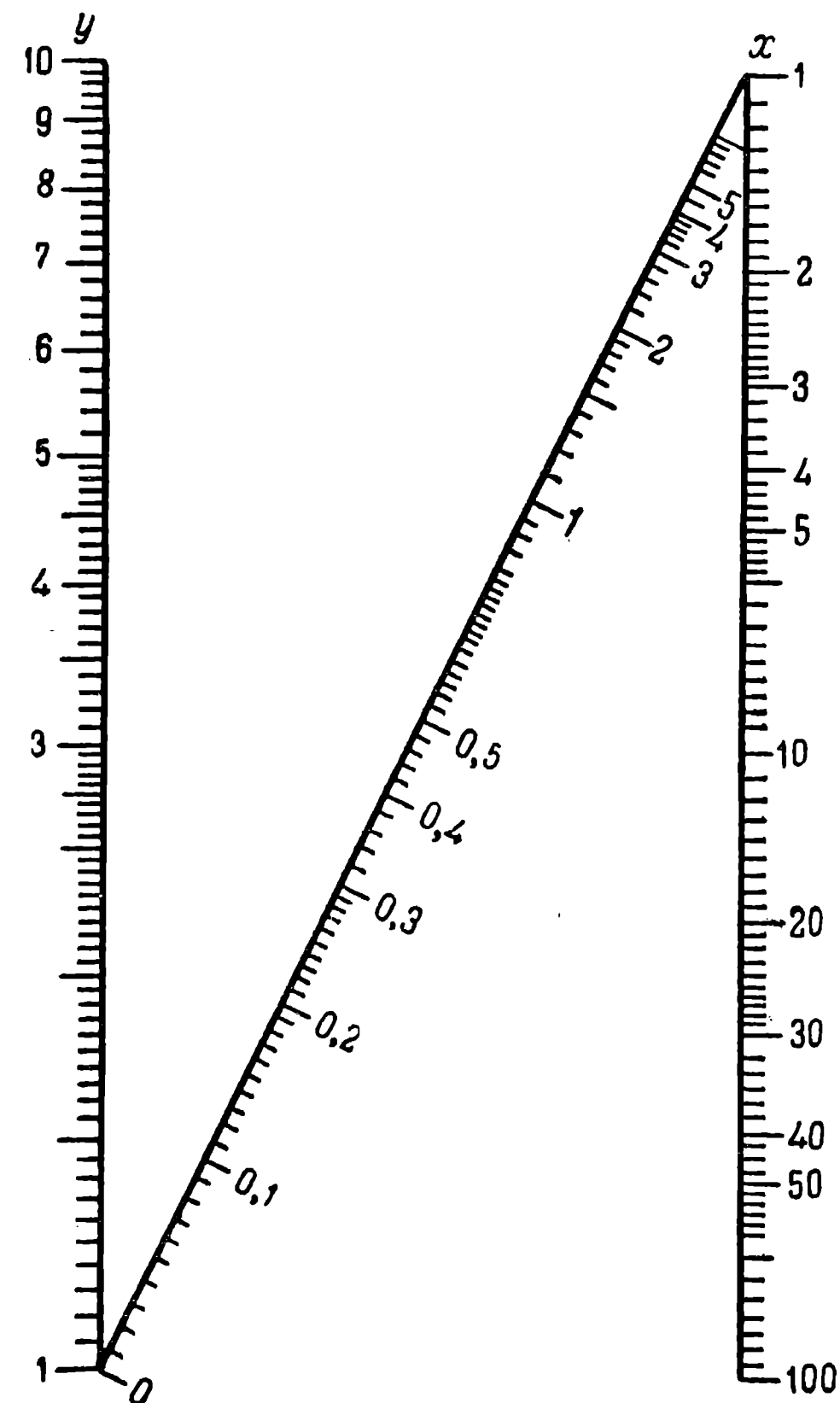


Рис. 44.

причём длину шкал берём равной попрежнему 20 см, а расстояние между ними равным, например, 10 см, как это

взято на рис. 44, который дан в масштабе около 1:2 по отношению к натуральной величине. Модули m_1 и m_2 определяем, исходя из общей длины этих шкал и пределов изменения y и x . Находим такие значения модулей:

$$m_1 = \frac{200 \text{ мм}}{\lg 10 - \lg 1} = 200 \text{ мм},$$

$$m_2 = \frac{200 \text{ мм}}{\lg 100 - \lg 1} = 100 \text{ мм}.$$

Третью шкалу строим на основании уравнения (15), определив предварительно значение h . Нулевые точки параллельных шкал лежат в пределах номограммы, поэтому определение h и построение наклонной шкалы не вызывают затруднений.

ГЛАВА VIII ОСНОВНЫЕ ПРИЁМЫ ИНТЕРПОЛИРОВАНИЯ

1. Линейное и графическое интерполирование. Если две физические величины зависят одна от другой, то мы обыкновенно измеряем значения одной величины при определённых значениях другой. Так, например, мы измеряем удельное сопротивление тел при определённых температурах, или угол отклонения в диффракционном спектре лучей определённой длины волны, или радиоактивную силу препарата в определённые моменты времени и т. п.

Во всех таких случаях можно сказать, что мы определяем значение некоторой функции

$$y = f(x). \quad (1)$$

при определённых значениях её аргумента x . В результате мы получаем ряд значений функции $y_0, y_1, y_2, y_3, \dots, y_m$, отвечающих значениям аргумента $x_0, x_1, x_2, x_3, \dots, x_m$. Часто вид функциональной зависимости f может оставаться при этом неизвестным, иногда функция может быть настолько сложной, что непосредственные вычисления с ней являются затруднительными. В обоих этих случаях принято заменять функцию $f(x)$ некоторой приближённой функцией $\varphi(x)$ возможно более простого вида, на которую мы накладываем непременное условие: при значениях аргумента, равных $x_0, x_1, x_2, x_3, \dots, x_m$, функция $\varphi(x)$ должна получать те же значения $y_0, y_1, y_2, y_3, \dots, y_m$, что и функция $f(x)$.

Поэтому можно написать:

$$\begin{aligned} y_0 &= f(x_0) = \varphi(x_0), \\ y_1 &= f(x_1) = \varphi(x_1), \\ y_2 &= f(x_2) = \varphi(x_2), \\ y_3 &= f(x_3) = \varphi(x_3), \\ &\dots \dots \dots \\ y_m &= f(x_m) = \varphi(x_m). \end{aligned} \quad (2)$$

Такую замену или представление функции $f(x)$ при помощи более простой приближённой функции $\varphi(x)$, которая удовлетворяет условиям (2), и можно назвать *интерполяцией*; соответственно этому функцию $\varphi(x)$ называют *интерполирующей* или *аппроксимирующей* функцией. Более узкой задачей интерполирования является определение значений функции y

при различных значениях аргумента x , которые лежат в интервалах между его значениями, данными непосредственно, т. е. при любых значениях x в интервале между x_0 и x_m .



Рис. 45.

Интерполированием в его наиболее простом виде мы пользуемся очень часто, иногда даже не зная этого. Так, например, если при взвешивании какого-либо тела мы применяем аналитические весы с достаточно большой чувствительностью,

то окончательную нагрузку весов, которая приводит их к равновесию, мы, как известно, не определяем из опыта, а находим из вычислений на основании определённой формулы. Действительно, при точном взвешивании мы непосредственно определяем только нагрузку весов, несколько меньшую веса тела, которая приводит их лишь к приближённому равновесию, так что соответствующая установка e_1 весов лежит близко от их нулевой точки e_0 , несколько справа от неё (рис. 45). Затем мы увеличиваем нагрузку весов обыкновенно на 1 мг и находим их новую установку e_2 , также недалёкую от нулевой точки e_0 , но по другую сторону от неё. Из этих результатов мы вычисляем те доли миллиграмма, которые следовало

бы добавить к общей нагрузке весов для того, чтобы их стрелка перешла от установки e_1 к нулевой точке e_0 . Это вычисление, для которого мы пользуемся формулой

$$p = \frac{e_1 - e_0}{e_1 - e_2} \text{ мг},$$

есть не что иное, как интерполирование в наиболее простой форме.

Точно так же, если мы делаем какое-либо вычисление при помощи логарифмов и нам необходимо найти логарифм такого числа, которое в таблицах логарифмов непосредственно не содержится, если, например, необходимо найти $\lg 147,26$ или $\lg \sin 41^\circ 32' 28''$, то мы всегда находим эти логарифмы при помощи простых вычислений; берём из таблиц логарифмы чисел, ближайших к данным, т. е. чисел 147,2 или $\sin 41^\circ 32'$, и затем, пользуясь таблицами пропорциональных частей, вносим в мантиссы этих логарифмов ту поправку, при которой мы получаем логарифмы чисел 147,26 или $\sin 41^\circ 32' 28''$. Эти общеизвестные приёмы, применяемые при нахождении чисел или логарифмов, являются также наиболее простым видом интерполирования.

При таком интерполировании мы предполагаем, что между данными величинами имеет место прямая пропорциональность, или, иначе говоря, что зависимость между ними графически выражается прямой линией. Вследствие этого такой вид интерполирования принято называть *линейным интерполированием*.

Линейное интерполирование, очевидно, можно применять: во-первых, в тех случаях, когда зависимость между y и x выражается или прямой линией, или кривой, но настолько близкой к прямой линии, что линейное интерполирование не угрожает грубыми ошибками, и, во-вторых, в тех случаях, когда имеется более сложная зависимость между y и x , выражаемая более сложной кривой, но мы интерполируем в пределах настолько малого изменения x , что на этом участке можно допустить без заметных погрешностей линейную зависимость между переменными.

Чрезвычайно часто применяется, далее, так называемое *графическое интерполирование*,

т. е. точное графическое воспроизведение результатов измерений, о чём мы подробно говорили раньше (гл. VI). Если на основании результатов измерений мы составляем точный график, то неизвестную нам функцию f мы заменяем её приближённым графическим изображением и определяем по графику значения y при каких угодно значениях x , если они лежат в пределах измерений. Такой графический метод определения приближённой, или интерполирующей, функции и принято называть графическим интерполированием.

Графическим интерполированием можно пользоваться при какой угодно зависимости между величинами, однако необходимо иметь в виду, что точное графическое построение сложных кривых иногда может оказаться очень затруднительным. Так, например, если кривая имеет резко выраженные максимумы или минимумы, положение которых не отвечает ни одной паре значений x и y , то в этих точках вычертить кривую можно только очень приближённо. Вследствие этого результаты, получаемые при помощи графического интерполирования, часто имеют очень ограниченную точность.

2. Интерполирование при равноотстоящих значениях аргумента. Таким образом, мы видим, что очень часто могут встречаться такие случаи, при которых невозможно применить ни линейного интерполирования вследствие сложной функциональной зависимости между x и y , ни графического интерполирования вследствие его ограниченной точности. Поэтому был выведен ряд различных интерполяционных формул; они дают выражение приближённой, или интерполирующей, функции $\varphi(x)$ и могут служить для определения значений y с большой степенью точности при любом виде функции $f(x)$, если она удовлетворяет условиям непрерывности.

При выводе интерполяционных формул обыкновенно предполагают, что интерполирующая функция $\varphi(x)$ может быть представлена в виде многочлена n -й степени по отношению к x , т. е. что можно положить:

$$y \approx \varphi(x) = a_0 x^n + a_1 x^{n-1} + a_2 x^{n-2} + a_3 x^{n-3} + \dots + a_n, \quad (3)$$

где через $a_0, a_1, a_2, a_3, \dots, a_n$ обозначены коэффициенты многочлена, которые являются неизвестными. Такой многочлен графически в системе координат (x, y) изображается кривой n -го порядка параболического типа, вследствие чего интерполяция этого вида получила название **параболической**.

Для того чтобы найти функцию $\varphi(x)$, необходимо, очевидно, определить степень многочлена n и величину его коэффициентов a . Рассуждая чисто теоретически, нетрудно показать, что определение коэффициентов a можно привести к решению системы линейных уравнений. Действительно, предполагая попрежнему, что нам дан ряд соответственных значений x и y :

$$\begin{aligned} x_0, x_1, x_2, x_3, \dots, x_m, \\ y_0, y_1, y_2, y_3, \dots, y_m, \end{aligned} \quad (4)$$

и подставляя эти значения последовательно в формулу (3), мы получим систему $(m+1)$ уравнений, в которых неизвестными служат $(n+1)$ коэффициентов $a_0, a_1, a_2, a_3, \dots, a_n$; по отношению к коэффициентам a эти уравнения являются линейными.

Однако на практике решение системы уравнений, даже линейных, при большом числе неизвестных может быть очень сложной вычислительной задачей; кроме того, остаётся неизвестной степень многочлена n .

Наконец, необходимо ещё иметь в виду, что задача об определении коэффициентов a становится неопределённой, если m меньше n , т. е. если число уравнений меньше числа неизвестных. Вполне определённое решение эта задача должна иметь в том случае, если m равно n . Первый случай, когда $m < n$, мы исключаем из рассмотрения, как крайне редко встречающийся на практике, и будем предполагать в дальнейшем, что m равно или больше n .

При выводе интерполяционных формул, весьма многочисленных и разнообразных, применяются различные, иногда достаточно сложные приёмы. Мы ограничимся рассмотрением только некоторых основных интерполяционных формул и изложим эти вопросы в возможно более элементарной форме.

Интерполирование оказывается наиболее простым в тех случаях, когда значения аргумента x образуют арифметическую прогрессию, т. е. когда интервалы между $x_0, x_1, x_2, \dots, x_m$ остаются постоянными, так что можно положить:

$$x_1 - x_0 = x_2 - x_1 = x_3 - x_2 = \dots = x_m - x_{m-1} = h,$$

где h — некоторая постоянная величина. Из этих равенств находим:

$$\begin{aligned} x_1 &= x_0 + h, \\ x_2 &= x_0 + 2h, \\ x_3 &= x_0 + 3h, \\ &\dots \end{aligned} \quad (5)$$

или в общем виде:

$$x_m = x_0 + mh. \quad (5')$$

Для дальнейшего следует обратить внимание на то, что в формуле (5') m является целым числом, одним из ряда натуральных чисел $0, 1, 2, 3, \dots$, однако эта формула оказывается справедливой и в том случае, если m получает дробные значения. Таким образом, опуская значок при x , формулу (5') можно написать ещё в таком виде:

$$x = x_0 + mh, \quad (5'')$$

где m может иметь какие угодно целые или дробные значения. Справедливость этого вывода следует из того, что если x рассматривать как функцию m и представить уравнение (5'') графически в системе координат x — ординаты и m — абсциссы, то получается прямая линия (рис. 46).

При выводе и применении на практике интерполяционных формул для случая равноотстоящих значений аргумента приходится оперировать с разностями значений y и разностями этих разностей, или, как принято говорить, с разностями функции $f(x)$ различных порядков. Если мы возьмём разности последовательных значений y , т. е. величины $y_1 - y_0, y_2 - y_1, y_3 - y_2, \dots, y_m - y_{m-1}$, то получаем разности первого порядка, или, короче, первые разности, которые принято обозначать

символом Δy , причём при y ставится индекс вычитаемого. Таким образом, можно написать:

$$\begin{aligned} \Delta y_0 &= y_1 - y_0, \\ \Delta y_1 &= y_2 - y_1, \\ \Delta y_2 &= y_3 - y_2, \\ &\dots \end{aligned} \quad (6)$$

Число первых разностей, очевидно, на единицу меньше числа значений y , т. е. равно $m - 1$.

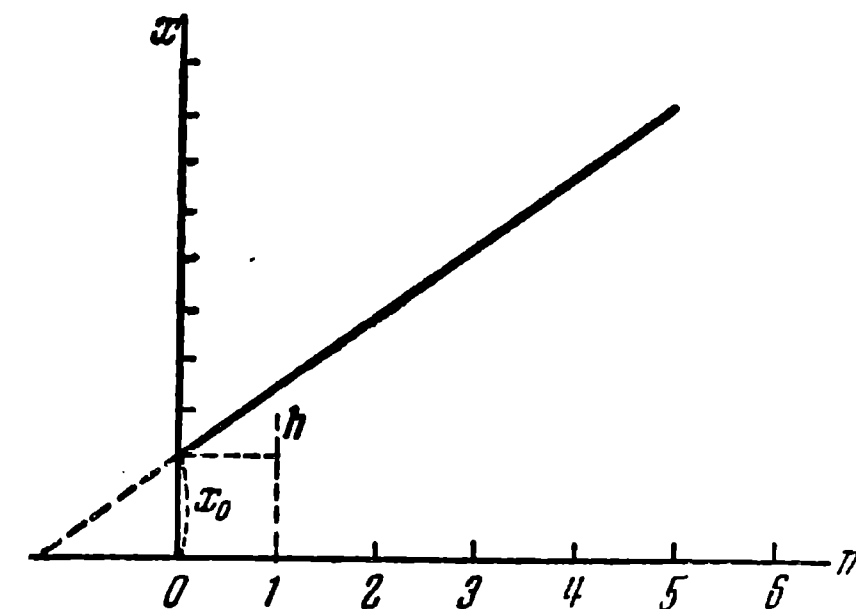


Рис. 46.

Разности первых разностей называются разностями второго порядка, или вторыми разностями, и обозначаются символом $\Delta^2 y$, причём при y ставится соответствующий индекс. Таким образом, принимая во внимание уравнения (6), можно написать

$$\begin{aligned} \Delta^2 y_0 &= \Delta y_1 - \Delta y_0 = y_2 - 2y_1 + y_0, \\ \Delta^2 y_1 &= \Delta y_2 - \Delta y_1 = y_3 - 2y_2 + y_1, \\ \Delta^2 y_2 &= \Delta y_3 - \Delta y_2 = y_4 - 2y_3 + y_2, \\ &\dots \end{aligned} \quad (6')$$

Число разностей второго порядка, очевидно, ещё на единицу меньше, т. е. равно $(m - 2)$. Далее, мы можем получить разности третьего порядка, или третьи разности, для которых, принимая во внимание уравнения (6) и (6'),

можно написать:

$$\begin{aligned}\Delta^3 y_0 &= \Delta^2 y_1 - \Delta^2 y_0 = y_3 - 3y_2 + 3y_1 - y_0, \\ \Delta^3 y_1 &= \Delta^2 y_2 - \Delta^2 y_1 = y_4 - 3y_3 + 3y_2 - y_1, \\ \Delta^3 y_2 &= \Delta^2 y_3 - \Delta^2 y_2 = y_5 - 3y_4 + 3y_3 - y_2, \\ &\dots\dots\dots\end{aligned}\quad (6'')$$

Число разностей третьего порядка равно, очевидно, $(m-3)$.

Точно так же мы можем получать выражения разностей следующих порядков, т. е. четвёртого, пятого и т. д., или, короче, четвёртые, пятые и т. д. разности.

Сравнивая правые части равенств (6), (6') и (6''), нетрудно установить закон, по которому составляются выражения разностей различных порядков. Действительно, мы видим, что в правых частях этих равенств: 1) индекс при y в каждом последующем члене понижается на единицу и 2) коэффициенты последовательных членов соответствуют коэффициентам разложения разности двух величин $(a-b)$ по формуле бинома Ньютона, причём степень бинома соответствует порядку разности.

При вычислении разностей на практике их принято располагать в виде таблицы, в которой выписываются также значения x и y ; разности в таблице обыкновенно пишут против промежутков тех чисел предыдущего столбца таблицы, из которых получается данная разность, как это показано в таблице XIII для разностей первых шести порядков. При таком расположении разностей их таблицу принято называть диагональной таблицей разностей.

Так как функция y приближённо выражается формулой (3) и представляет собой многочлен n -й степени относительно x , то нам необходимо найти выражения разностей различных порядков для подобных многочленов. Вычисления показывают, что последовательные разности многочлена n -й степени являются также многочленами, степени которых относительно x понижаются на единицу при переходе к следующей разности, т. е. при увеличении порядка разности на единицу. Так, первая разность

Таблица XIII

x	y	Δy	$\Delta^2 y$	$\Delta^3 y$	$\Delta^4 y$	$\Delta^5 y$	$\Delta^6 y$
x_0	y_0	Δy_0					
x_1	y_1	Δy_1	$\Delta^2 y_0$	$\Delta^3 y_0$			
x_2	y_2	Δy_2	$\Delta^2 y_1$	$\Delta^3 y_1$	$\Delta^4 y_0$	$\Delta^5 y_0$	$\Delta^6 y_0$
x_3	y_3	Δy_3	$\Delta^2 y_2$	$\Delta^3 y_2$	$\Delta^4 y_1$	$\Delta^5 y_1$	
x_4	y_4	Δy_4	$\Delta^2 y_3$	$\Delta^3 y_3$	$\Delta^4 y_2$		
x_5	y_5	Δy_5	$\Delta^2 y_4$				
x_6	y_6						

многочлена n -й степени, т. е.

$$\Delta \varphi(x) = \varphi(x+h) - \varphi(x),$$

является многочленом $(n-1)$ -й степени относительно x и может быть в общем виде представлена таким выражением:

$$\Delta \varphi(x) = b_0 x^{n-1} + b_1 x^{n-2} + b_2 x^{n-3} + \dots + b_{n-2} x + b_{n-1}.$$

Точно так же вторую разность многочлена n -й степени в общем виде можно представить таким многочленом:

$$\Delta^2 \varphi(x) = c_0 x^{n-2} + c_1 x^{n-3} + c_2 x^{n-4} + \dots + c_{n-3} x + c_{n-2}.$$

Соответственно этому последующие разности многочлена n -й степени могут быть представлены в общем виде выражениями

$$\Delta^3 \varphi(x) = d_0 x^{n-3} + d_1 x^{n-4} + d_2 x^{n-5} + \dots + d_{n-4} x + d_{n-3},$$

$$\Delta^4 \varphi(x) = e_0 x^{n-4} + e_1 x^{n-5} + e_2 x^{n-6} + \dots + e_{n-5} x + e_{n-4}.$$

Продолжая вычисления последовательных разностей всё более высокого порядка, мы приходим к выводу, что разность n -го порядка многочлена n -й степени содержит x в нулевой степени, т. е. является постоянным числом. Отсюда следует также, что все разности порядка выше n для многочлена n -й степени оказываются

равными нулю. Все эти выводы являются вполне аналогичными тем, которые доказываются в дифференциальном исчислении по отношению к производным различных порядков от функции n -й степени относительно x .

Вывод, согласно которому разность n -го порядка для многочлена n -й степени оказывается постоянным числом, является весьма важным для определения степени n интерполирующего многочлена, что становится понятным из следующих соображений. Если нам дан ряд (4) значений x и y и, вычисляя таблицу разностей, мы находим, что разности некоторого порядка, например порядка k , оказываются постоянными, то можно утверждать, что функциональная зависимость y от x в данном случае может быть представлена в виде многочлена также степени k . Таким образом, мы видим, что вычисление последовательных разностей даёт нам возможность определить степень того многочлена, который может служить интерполирующей функцией (x) в данном случае.

Пользуясь этим приёмом определения степени интерполирующей функции, необходимо вместе с тем иметь в виду следующее:

1) Эти выводы можно применять только в тех случаях, когда величина h в выражениях (5) или (5'') остаётся постоянной.

2) Обработывая результаты физических измерений, мы оперируем с числами не точными, а приближёнными. Поэтому, составляя таблицу разностей для результатов физических измерений, мы не вправе ожидать, что во всех случаях разности некоторого порядка окажутся действительно строго постоянными числами. На практике приходится довольствоваться обыкновенно лишь относительным постоянством разностей, которые могут несколько отличаться одна от другой. Однако это различие не должно выходить далеко за пределы тех десятичных знаков, которые в результатах измерений являются сомнительными, т. е. содержат некоторые ошибки.

Необходимо сделать ещё одно очень существенное замечание. Если при измерении величин y или при вычислении их значений мы допустили какую-либо ошибку или если эта ошибка появилась при вычислении разностей,

то это обстоятельство немедленно находит отражение в таблице разностей: оказывается, что допущенная нами ошибка начинает охватывать всё более широкую зону разностей и постепенно возрастает. Действительно, допустим, что значение y_4 мы ввели в таблицу с некоторой ошибкой, равной ε , т. е. вместо правильного значения y_4 в таблице оказалось его ошибочное значение $y_4 + \varepsilon$. В таком случае, как мы видим из таблицы XIV, при вычислении первых разностей ошибка ε переходит в две соседние строки таблицы, при вычислении вторых разностей она распространяется на три строки и в средней строке возрастает в два раза, затем ошибка охватывает четыре строки и возрастает в три раза и т. д., причём все эти ошибки располагаются симметрично около той строки, где впервые была допущена ошибка.

Таблица XIV

x	y	Δy	$\Delta^2 y$	$\Delta^3 y$	$\Delta^4 y$
x_0	y_0	Δy_0			
x_1	y_1	Δy_1	$\Delta^2 y_0$		
x_2	y_2	Δy_2	$\Delta^2 y_1$	$\Delta^3 y_0$	
x_3	y_3	$\Delta y_3 + \varepsilon$	$\Delta^2 y_2 + \varepsilon$	$\Delta^3 y_1 + \varepsilon$	$\Delta^4 y_0 + \varepsilon$
x_4	$y_4 + \varepsilon$	$\Delta y_4 - \varepsilon$	$\Delta^2 y_3 - 2\varepsilon$	$\Delta^3 y_2 - 3\varepsilon$	$\Delta^4 y_1 - 4\varepsilon$
x_5	y_5	Δy_5	$\Delta^2 y_4 + \varepsilon$	$\Delta^3 y_3 + 3\varepsilon$	$\Delta^4 y_2 + 6\varepsilon$
x_6	y_6	Δy_6	$\Delta^2 y_5$	$\Delta^3 y_4 - \varepsilon$	$\Delta^4 y_3 - \varepsilon$
x_7	y_7	Δy_7	$\Delta^2 y_6$	$\Delta^3 y_5$	$\Delta^4 y_4 + \varepsilon$
x_8	y_8	Δy_8	$\Delta^2 y_7$	$\Delta^3 y_6$	$\Delta^4 y_5$
x_9	y_9				

Отсюда следует, что если, вычисляя последовательные разности, мы не наблюдаем для них даже относительного постоянства, и отклонения разностей от их среднего арифметического начинают систематически возрастать, то есть основание предполагать наличие ошибки в вычислениях или измерениях. В таких случаях необходимо прежде

всего, проверить вычисления. Если в них ошибки не будет обнаружено, то необходимо проверить те измерения y , которые лежат в областях, где отклонения ошибок особенно велики. Если и здесь не будет обнаружено никаких дефектов, то объяснить отсутствие постоянных разностей можно только особенностями данного физического явления, непрерывность которого в некоторых областях может оказаться нарушенной (стр. 203).

Рассмотрим несколько примеров.

Пример 1. При калибровании пружинных весов их нагрузку P постепенно увеличивали на 2 г, отсчитывая каждый раз показания N весов; результаты отсчётов приведены в таблице XV.

Таблица XV

P	0	2	4	6	8	10	12	14	16	18	20
N	17,30	18,24	19,18	20,12	21,06	22,00	22,94	23,88	24,82	25,76	26,70

Определим степень интерполирующего многочлена. Для этого составляем таблицу разностей (таблица XVI).

Таблица XVI

P	N	ΔN	$\Delta^2 N$	P	N	ΔN	$\Delta^2 N$
0	17,30			10	22,00		
2	18,24	0,94		12	22,94	0,94	0
4	19,18	0,94	0	14	23,88	0,94	0
6	20,12	0,94	0	16	24,82	0,94	0
8	21,06	0,94	0	18	25,76	0,94	0
10	22,00	0,94		20	26,70	0,94	

Мы видим, что уже первые разности оказываются постоянными. Таким образом, можно утверждать, что в данном случае интерполирующей функцией служит многочлен первой степени, т. е. показания весов оказываются пропорциональными их нагрузке.

Этот вывод вполне понятен, так как при небольших нагрузках пружинных весов деформация пружины следует закону упругих деформаций, согласно которому между величиной деформации и силой, её вызывающей, должна иметь место прямая пропорциональность.

Пример 2. Плотность D воды в области температур между 0° и 100° С была, как известно, измерена с большой степенью точности. Результаты этих измерений в интервале температур от 20 до 25° С для каждого градуса даны в таблице XVII.

Таблица XVII

t	20°	21°	22°	23°	24°	25°
D	0,998230	0,998019	0,997797	0,997565	0,997323	0,997071

Определим степень того многочлена, которым может быть представлена зависимость плотности воды от температуры в указанном интервале температур. Для этого составляем таблицу разностей (таблица XVIII).

Таблица XVIII

t	D	ΔD	$\Delta^2 D$	$\Delta^3 D$
20°	0,998230			
21°	8019	-0,000211		
22°	7797	222	-0,000011	
23°	7565	232	10	-0,000001
24°	7323	242	10	0
25°	7071	252	10	0

Все остальные члены многочлена (7) оказываются равными нулю. Из этого уравнения, принимая во внимание уравнение (8), находим:

$$c_1 = \frac{y_1 - y_0}{h}.$$

Отсюда на основании первого из уравнений (6) получаем:

$$c_1 = \frac{\Delta y_0}{h}. \quad (8')$$

3) При $x = x_2$ имеем соответственно:

$$\varphi(x_2) = y_2 = c_0 + c_1(x_2 - x_0) + c_2(x_2 - x_0)(x_2 - x_1).$$

Подставляя в это выражение значения c_0 и c_1 из уравнений (8) и (8'), получаем:

$$y_2 = y_0 + \frac{\Delta y_0}{h} \cdot 2h + c_2 \cdot 2h \cdot h = y_0 + \frac{y_1 - y_0}{h} \cdot 2h + c_2 \cdot 2h^2.$$

Решая это уравнение относительно c_2 , находим:

$$c_2 = \frac{(y_2 - 2y_1 + y_0)}{2h^2}.$$

Отсюда на основании первого из уравнений (6') получаем:

$$c_2 = \frac{\Delta^2 y_0}{2h^2}. \quad (8'')$$

4) При $x = x_3$

$$\varphi(x_3) = y_3 = c_0 + c_1(x_3 - x_0) + c_2(x_3 - x_0)(x_3 - x_1) + c_3(x_3 - x_0)(x_3 - x_1)(x_3 - x_2)$$

или

$$\begin{aligned} y_3 &= c_0 + c_1 3h + c_2 \cdot 3h \cdot 2h + c_3 \cdot 3h \cdot 2h \cdot h = \\ &= y_0 + \frac{y_1 - y_0}{h} \cdot 3h + \frac{y_2 - 2y_1 + y_0}{2h^2} \cdot 3h \cdot 2h + c_3 \cdot 2 \cdot 3h^3. \end{aligned}$$

Решая это уравнение относительно c_3 , находим:

$$c_3 = \frac{y_3 - 3y_2 + 3y_1 - y_0}{2 \cdot 3 \cdot h^3}.$$

Отсюда на основании первого из уравнений (6'') получаем:

$$c_3 = \frac{\Delta^3 y_0}{2 \cdot 3 \cdot h^3}. \quad (8''')$$

5) При $x = x_4$

$$\begin{aligned} \varphi(x_4) = y_4 &= c_0 + c_1(x_4 - x_0) + \\ &+ c_2(x_4 - x_0)(x_4 - x_1) + c_3(x_4 - x_0)(x_4 - x_1)(x_4 - x_2) + \\ &+ c_4(x_4 - x_0)(x_4 - x_1)(x_4 - x_2)(x_4 - x_3). \end{aligned}$$

Применяя к этому уравнению преобразования, вполне аналогичные тем, которыми мы пользовались в предыдущих случаях, мы приводим его к виду

$$y_4 = -y_0 + 4y_1 - 6y_2 + 4y_3 + c_4 \cdot 2 \cdot 3 \cdot 4 \cdot h^4.$$

Решая это уравнение относительно c_4 и принимая во внимание уравнения для $\Delta^4 y$, аналогичные уравнениям (6''), получаем:

$$c_4 = \frac{\Delta^4 y_0}{2 \cdot 3 \cdot 4 \cdot h^4}. \quad (8''')$$

Из выражений, полученных для первых пяти коэффициентов многочлена (7), нетрудно установить общую формулу для определения всех коэффициентов этого многочлена, начиная с c_1 :

$$c_m = \frac{\Delta^m y_0}{m! h^m}, \quad (9)$$

где m может иметь все значения ряда натуральных чисел, от единицы до n .

Подставляя полученные значения коэффициентов $c_0, c_1, c_2, c_3, \dots$ в выражение (7), находим:

$$\begin{aligned} y \approx \varphi(x) &= y_0 + \frac{\Delta y_0}{h} (x - x_0) + \frac{\Delta^2 y_0}{2! h^2} (x - x_0)(x - x_1) + \\ &+ \frac{\Delta^3 y_0}{3! h^3} (x - x_0)(x - x_1)(x - x_2) + \dots \\ &\dots + \frac{\Delta^n y_0}{n! h^n} (x - x_0)(x - x_1)(x - x_2)(x - x_3) \dots (x - x_{n-1}). \end{aligned} \quad (10)$$

В формуле (10) значения коэффициентов выражаются через известные величины и могут быть вычислены;

Рассмотрим несколько примеров применения формулы Ньютона.

Пример 1. Исследование зависимости двух величин x и y показало, что при постепенном увеличении значений x на единицу y получает значения, указанные в таблице XXI.

Таблица XXI

x	1	2	3	4	5	6	7
y	3	7	13	21	31	43	57

Найдём степень интерполирующего многочлена и вычислим значения y , отвечающие каким-либо промежуточным значениям, например 3,1 и 4,67.

Прежде всего необходимо составить таблицу разностей:

Таблица XXII

x	y	Δy	$\Delta^2 y$	$\Delta^3 y$
1	3	4		
2	7	6	2	0
3	13	8	2	0
4	21	10	2	0
5	31	12	2	0
6	43	14	2	
7	57			

Мы видим, что вторые разности остаются постоянными, таким образом, можно утверждать, что интерполирующую функцию в данном случае можно представить многочленом второй степени.

Вычисляем значения y при x , равном 3,1 и 4,67.

1) При $x=3,1$ мы должны интерполировать в промежутке между значениями x , равными 3 и 4, поэтому полагаем:

$$x_0 = 3, x_1 = 4.$$

Соответственные значения y и его разностей берём из таблицы XXII:

$$y_0 = 13; y_1 = 21; \Delta y_0 = 8; \Delta^2 y_0 = 2.$$

Величина h в данном случае равна единице:

$$h = 1.$$

Из формулы (5*) находим значение n :

$$n = \frac{x - x_0}{h} = \frac{3,1 - 3}{1} = 0,1.$$

В формулах (12) в данном случае мы берём первые три члена, так как интерполирующей функцией служит многочлен второй степени. Таким образом, первая из формул (12) в данном случае получает вид

$$y = y_0 + \frac{\Delta y_0}{h} (x - x_0) + \frac{\Delta^2 y_0}{2h^2} (x - x_0) (x - x_1).$$

Переходя в правой части к числовым значениям, находим:

$$y = 13 + \frac{8}{1} (3,1 - 3) + \frac{2}{2 \cdot 1} (3,1 - 3) (3,1 - 4)$$

или

$$y = 13 + 8 \cdot 0,1 - 0,1 \cdot 0,9 = 13,71.$$

Точно так же можно применить вторую из формул (12), которая в данном случае получает вид

$$y = y_0 + n \Delta y_0 + \frac{n(n-1)}{2!} \Delta^2 y_0.$$

Подставляя в правую часть числовые значения, получаем:

$$y = 13 + 0,1 \cdot 8 - \frac{0,1 \cdot 0,9}{2!} \cdot 2 = 13,71,$$

т. е. тот же самый результат.

2) При $x=4,67$ мы должны интерполировать в промежутке между значениями x , равными 4 и 5. Поэтому

аналогично первому случаю имеем:

$$\begin{aligned}x_0 &= 4, & y_0 &= 21, \\x_1 &= 5, & y_1 &= 31.\end{aligned}$$

Находим значение n :

$$n = \frac{x - x_0}{h} = \frac{4,67 - 4}{1} = 0,67.$$

Применяя вторую из формул (12), получаем:

$$y = y_0 + n\Delta y_0 + \frac{n(n-1)}{2!} \Delta^2 y_0.$$

Подставляем в эту формулу числовые значения:

$$y = 21 + 0,67 \cdot 10 + \frac{0,67(0,67-1)}{2} \cdot 2,$$

или

$$y = 21 + 0,67 \cdot 10 - 0,67 \cdot 0,33 = 27,48.$$

При вычислении значений y мы брали в качестве x_0 и x_1 два значения x , ближайшие к тому, для которого вычисляется значение y ; так, при $x = 3,1$ мы считали x_0 и x_1 равными 3 и 4. Но это не является обязательным, и мы можем считать x_0 равным любому из его значений, которые в таблице стоят выше данного значения x ; так, при $x = 3,1$ мы можем считать x_0 равным 1, или 2, или 3. Результаты вычисления во всех случаях остаются одинаковыми. Действительно, примем, например, что

$$x_0 = 1 \text{ и } x_1 = 2.$$

Соответственно этому имеем:

$$y_0 = 3; \quad y_1 = 7; \quad \Delta y_0 = 4 \text{ и } \Delta^2 y_0 = 2.$$

Отсюда находим:

$$n = \frac{3,1 - 1}{h} = 2,1.$$

Вследствие этого, применяя первую из формул (12), находим:

$$\begin{aligned}y &= 3 + \frac{4}{1} (3,1 - 1) + \frac{2}{2 \cdot 1} (3,1 - 1) (3,1 - 2) = \\&= 3 + 4 \cdot 2,1 + 2,1 \cdot 1,1 = 13,71.\end{aligned}$$

Вторая из формул (12) в этом случае даёт:

$$y = 3 + 2,1 \cdot 4 + \frac{2,1 \cdot 1,1}{2} \cdot 2 = 13,71.$$

Точно так же, если бы мы положили, что

$$x_0 = 2 \text{ и } x_1 = 3,$$

то имели бы:

$$y_0 = 7; \quad y_1 = 13; \quad \Delta y_0 = 6 \text{ и } \Delta^2 y_0 = 2.$$

Для n в этом случае получаем:

$$n = \frac{3,1 - 2}{h} = 1,1.$$

Таким образом, применяя первую и вторую формулы (12), находим:

$$\begin{aligned}y &= 7 + \frac{6}{1} (3,1 - 2) + \frac{2}{2} (3,1 - 2) (3,1 - 3) = \\&= 7 + 6 \cdot 1,1 + 1,1 \cdot 0,1 = 13,71, \\y &= 7 + 1,1 \cdot 6 + \frac{1,1 \cdot 0,1 \cdot 2}{2} = 13,71.\end{aligned}$$

Отсюда следует, что мы можем брать различные значения x , что даёт возможность проверять результаты вычислений.

Те же формулы (12) дают возможность определить коэффициенты того многочлена второй степени, которым, как мы установили, выражается в данном случае зависимость между y и x . Для этого воспользуемся первой из формул (12). В этой формуле раскроем скобки в правой части и затем расположим многочлен по нисходящим степеням x .

Находим:

$$\begin{aligned}y &= \frac{\Delta^2 y_0}{2h^2} x^2 + \left(\frac{\Delta y_0}{h} - x_0 \frac{\Delta^2 y_0}{2h^2} - x_1 \frac{\Delta^2 y_0}{2h^2} \right) x + \\&\quad + \left(y_0 - x_0 \frac{\Delta y_0}{h} + x_0 x_1 \frac{\Delta^2 y_0}{2h^2} \right).\end{aligned}$$

Вычисляем числовые значения коэффициентов при x и свободного члена в правой части этой формулы.

Получаем:

$$\frac{\Delta^2 y_0}{2h^2} = \frac{2}{2 \cdot 1} = 1,$$

$$\frac{\Delta y_0}{h} - x_0 \frac{\Delta^2 y_0}{2h^2} - x_1 \frac{\Delta^2 y_0}{2h^2} = \frac{8}{1} - 3 \cdot \frac{2}{2} - 4 \cdot \frac{2}{2 \cdot 1} = 1,$$

$$y_0 - x_0 \frac{\Delta y_0}{h} + x_0 x_1 \frac{\Delta^2 y_0}{2h^2} = 13 - 3 \cdot \frac{8}{1} + 3 \cdot 4 \cdot \frac{2}{2} = 1.$$

Оба коэффициента при x и свободный член оказываются равными единице, таким образом, в данном случае функция y может быть представлена таким многочленом:

$$y = x^2 + x + 1. \quad (13)$$

Нетрудно непосредственным вычислением убедиться в правильности этого результата, а также в правильности найденных выше значений y . Весьма простое решение всех вопросов в данном примере объясняется тем, что числа в таблице XXI являются не приближёнными, а точными, и все значения y непосредственно вычисляются из уравнения (13).

Пример 2. При исследовании некоторой химической реакции через каждые 5 минут определялось количество A вещества, оставшееся в системе. Результаты измерений приведены в таблице XXIII, где даны: время t после начала реакции в минутах и количество A вещества в процентах.

Таблица XXIII

t	7	12	17	22	27	32	37
A	83,7	72,9	63,2	54,7	47,5	41,4	36,3

Определим, какой процент вещества остаётся в системе по истечении, например, 25 минут после начала реакции.

Составляем таблицу разностей; результаты даны в таблице XXIV.

Таблица XXIV

t	A	ΔA	$\Delta^2 A$	$\Delta^3 A$	$\Delta^4 A$
7	83,7				
12	72,9	-10,8			
17	63,2	-9,7	1,1		
22	54,7	-8,5	1,2	0,1	
27	47,5	-7,2	1,3	0,1	0
32	41,4	-6,1	1,1	-0,2	-0,3
37	36,3	-5,1	1,0	-0,1	+0,1

Из таблицы разностей мы видим, что была, повидимому, допущена некоторая ошибка при t , равном 22 минутам, но тем не менее вторые разности можно считать практически постоянными, несмотря на небольшие отклонения их от среднего значения, которое примем равным 1,1. Поэтому, взяв для вычислений первую из формул (12), ограничиваемся в ней первыми тремя членами:

$$y = y_0 + \frac{\Delta y_0}{h}(x - x_0) + \frac{\Delta^2 y_0}{2h^2}(x - x_0)(x - x_1).$$

В данном случае имеем: $x = 25$. Соответственно этому принимаем:

$$x_0 = 22, \quad x_1 = 27, \\ y_0 = 54,7 \quad y_1 = 47,5.$$

Из таблицы разностей видим, что

$$h = 5 \text{ и } \Delta y = -7,2.$$

Для $\Delta^2 y_0$ берём среднее значение 1,1, т. е. полагаем:

$$\Delta^2 y_0 = 1,1.$$

Вставляя эти значения в интерполяционную формулу, находим:

$$y = 54,7 - \frac{7,2}{5}(25 - 22) + \frac{1,1}{2 \cdot 5^2}(25 - 22)(25 - 27)$$

или

$$y = 54,7 - \frac{7,2 \cdot 3}{5} - \frac{1,1 \cdot 3 \cdot 2}{2 \cdot 25}.$$

Вычисляя, находим:

$$y = 50,25,$$

т. е. через 25 минут после начала реакции в системе остаётся приблизительно 50,2% начального количества вещества.

Пример 3. Воспользуемся значениями плотности воды при разных температурах, приведённых в таблице XVII, и вычислим плотность воды при какой-либо промежуточной температуре, например при 22,7°С. Так как в данном случае вторые разности оказываются постоянными, то попрежнему можем ограничиться тремя первыми членами в формулах (12); таким образом, применяя вторую из этих формул, пишем её в таком виде:

$$y = y_0 + n\Delta y_0 + \frac{n(n-1)}{1 \cdot 2} \Delta^2 y_0.$$

В данном случае мы интерполируем в промежутке между 22 и 23°С, поэтому принимаем, что

$$x_0 = 22; y_0 = 0,997797.$$

Из таблицы разностей мы видим, что

$$h = 1; \Delta y = -0,000232; \Delta^2 y = -0,000010.$$

По формуле (5*) находим значение n :

$$n = \frac{x - x_0}{h} = \frac{22,7 - 22,0}{1} = 0,7.$$

Подставляем числовые значения в интерполяционную формулу:

$$y = 0,997797 - 0,7 \cdot 0,000232 + \frac{0,7 \cdot 0,3}{2} \cdot 0,000010$$

или

$$y = 0,997797 - 0,7 \cdot 0,000232 + 0,21 \cdot 0,000005.$$

Отсюда находим:

$$y = 0,997636.$$

Таким образом, плотность воды при температуре 22,7°С оказывается равной 0,997636 см⁻³г.

Пример 4. В таблице XXV приведены значения эллиптического интеграла первого рода, определяемого

Таблица XXV

α	K	ΔK	$\Delta^2 K$	$\Delta^3 K$	$\Delta^4 K$	$\Delta^5 K$	$\Delta^6 K$
75°	2,76806	0,06461	0,00528	0,00084	0,00019	0,00013	-0,00005
76°	2,83267						
77°	2,90256	06989	0612	103	032	08	18
78°	2,97857	07601	0715				
79°	3,06173	08316	0850	135	040	26	1
80°	3,15339	09166	1025	175	066	25	43
81°	3,25530	10191	1266	241	091	68	
82°	3,36987	11457	1598	332	159		
83°	3,50042	13055	2089	491			
84°	3,65186	15144					

формулой

$$K = F(\alpha) = \int_0^{\frac{\pi}{2}} \frac{d\varphi}{\sqrt{1 - \sin^2 \alpha \cdot \sin^2 \varphi}}.$$

Эти значения даны с пятью десятичными знаками для углов α , выражаемых целым числом градусов, в интервале от 75 до 84°.

Найдём значение K для какого-либо промежуточного значения α , например для $78^\circ 30'$.

Определяя разности различных порядков, мы видим, что в данном случае только пятые разности можно принять условно постоянными, считая, что резко выраженный неправильный ход шестых разностей объясняется теми ошибками, которые неизбежно появляются при округлении значений K . Таким образом, принимая пятые разности постоянными, находим их среднее значение; оно оказывается равным 0,00028.

Для вычислений K воспользуемся второй из формул (12), полагая

$$\begin{aligned}x_0 &= 78^\circ, & y_0 &= 2,97857, \\x_1 &= 79^\circ, & y_1 &= 3,06173.\end{aligned}$$

Кроме того, имеем:

$$x = 78^\circ 30' = 78,5^\circ; \quad h = 1 \quad \text{и} \quad n = \frac{x - x_0}{h} = 0,5.$$

Вследствие этого, ограничиваясь во второй из формул (12) шестью членами и подставляя в неё числовые данные, находим:

$$\begin{aligned}K &= 2,97857 + 0,5 \cdot 0,08316 + \frac{0,5(0,5-1)}{2!} \cdot 0,00850 + \\&+ \frac{0,5 \cdot (0,5-1)(0,5-2)}{3!} \cdot 0,00175 + \\&+ \frac{0,5 \cdot (0,5-1)(0,5-2)(0,5-3)}{4!} \cdot 0,00066 + \\&+ \frac{0,5(0,5-1)(0,5-2)(0,5-3)(0,5-4)}{5!} \cdot 0,00028\end{aligned}$$

или

$$\begin{aligned}K &= 2,97857 + 0,04158 - \frac{0,5 \cdot 0,5}{2} \cdot 0,00850 + \\&+ \frac{0,5 \cdot 0,5 \cdot 1,5}{2 \cdot 3} \cdot 0,00175 - \frac{0,5 \cdot 0,5 \cdot 1,5 \cdot 2,5}{2 \cdot 3 \cdot 4} \cdot 0,00066 + \\&+ \frac{0,5 \cdot 0,5 \cdot 1,5 \cdot 2,5 \cdot 3,5}{2 \cdot 3 \cdot 4 \cdot 5} \cdot 0,00028.\end{aligned}$$

Производя вычисления, находим:

$$K = 3,019179,$$

т. е. с точностью до пятого десятичного знака:

$$K = 3,01918.$$

Этот пример представляет интерес в том отношении, что мы получили ряд с очень небольшой сходимостью, и для того, чтобы найти значение K с той точностью, с которой его значения даны в таблице, нам пришлось вычислить сумму первых шести членов этого ряда. Если бы мы, желая упростить вычисления, ограничились суммой пяти первых членов, то получили бы значение K меньше найденного на одну единицу последнего десятичного знака, в чём нетрудно убедиться непосредственным вычислением величины шестого члена.

4. Интерполирование при неравноотстоящих значениях аргумента. Если функция $f(x)$ задана попрежнему рядом значений x и y подобно ряду (4), но интервалы значений x оказываются различными, то применять формулы (12) не представляется возможным, так как они выведены на основании формул (5), которые в этом случае утрачивают своё значение. В этих случаях переходят к вычислению так называемых разделённых или приведённых разностей, которые мы получаем, если разность значений y или значения его разностей делим на разность значений x . Иначе говоря, мы вычисляем величину разностей, которые отвечают единице изменения x . Таким образом, сохраняя прежние символы, но с чёрточкой сверху, мы можем для определения приведённых разностей различных порядков составить такие выражения:

Приведённые разности первого порядка:

$$\bar{\Delta}y_0 = \frac{y_1 - y_0}{x_1 - x_0}; \quad \bar{\Delta}y_1 = \frac{y_2 - y_1}{x_2 - x_1}; \quad \bar{\Delta}y_2 = \frac{y_3 - y_2}{x_3 - x_2} \dots$$

Приведённые разности второго порядка:

$$\bar{\Delta}^2y_0 = \frac{\bar{\Delta}y_1 - \bar{\Delta}y_0}{x_2 - x_0}; \quad \bar{\Delta}^2y_1 = \frac{\bar{\Delta}y_2 - \bar{\Delta}y_1}{x_3 - x_1}; \quad \bar{\Delta}^2y_2 = \frac{\bar{\Delta}y_3 - \bar{\Delta}y_2}{x_4 - x_2} \dots$$

Приведённые разности третьего порядка:

$$\begin{aligned}\bar{\Delta}^3 y_0 &= \frac{\bar{\Delta}^2 y_1 - \bar{\Delta}^2 y_0}{x_3 - x_0}; \\ \bar{\Delta}^3 y_1 &= \frac{\bar{\Delta}^2 y_2 - \bar{\Delta}^2 y_1}{x_4 - x_1}; \quad \bar{\Delta}^3 y_2 = \frac{\bar{\Delta}^2 y_3 - \bar{\Delta}^2 y_2}{x_5 - x_2} \dots\end{aligned}$$

Аналогичными выражениями определяются приведённые разности и более высоких порядков, причём если функцию $f(x)$ мы приближённо представляем в виде многочлена степени m , то приведённая разность порядка m будет постоянной величиной, а приведённые разности более высоких порядков оказываются равными нулю, как это имело место и при равноотстоящих значениях аргумента.

Для приведённых разностей принято составлять таблицы, в простейших случаях совершенно аналогичные тем, которые мы имели раньше, как это показано в таблице XXVI, где даны приведённые разности первых пяти порядков.

Т а б л и ц а XXVI

x	y	$\bar{\Delta}y$	$\bar{\Delta}^2y$	$\bar{\Delta}^3y$	$\bar{\Delta}^4y$	$\bar{\Delta}^5y$
x_0	y_0	$\bar{\Delta}y_0$				
x_1	y_1	$\bar{\Delta}y_1$	$\bar{\Delta}^2y_0$	$\bar{\Delta}^3y_0$		
x_2	y_2	$\bar{\Delta}y_2$	$\bar{\Delta}^2y_1$	$\bar{\Delta}^3y_1$	$\bar{\Delta}^4y_0$	$\bar{\Delta}^5y_0$
x_3	y_3	$\bar{\Delta}y_3$	$\bar{\Delta}^2y_2$	$\bar{\Delta}^3y_2$	$\bar{\Delta}^4y_1$	$\bar{\Delta}^5y_1$
x_4	y_4	$\bar{\Delta}y_4$	$\bar{\Delta}^2y_3$	$\bar{\Delta}^3y_3$	$\bar{\Delta}^4y_2$	$\bar{\Delta}^5y_2$
x_5	y_5	$\bar{\Delta}y_5$	$\bar{\Delta}^2y_4$	$\bar{\Delta}^3y_4$	$\bar{\Delta}^4y_3$	
x_6	y_6	$\bar{\Delta}y_6$	$\bar{\Delta}^2y_5$			
x_7	y_7					

Интерполирующая функция $\varphi(x)$ определяется выражением, которое имеет такой вид:

$$f(x) \approx \varphi(x) = y_0 + (x - x_0) \bar{\Delta}y_0 + (x - x_0)(x - x_1) \bar{\Delta}^2y_0 + (x - x_0)(x - x_1)(x - x_2) \bar{\Delta}^3y_0 + \dots, \quad (14)$$

где $x_0, x_1, x_2, x_3, \dots$ обозначают последовательные значения x в таблице приведённых разностей.

Число членов формулы (14), которое мы вводим при интерполировании, определяется, так же как и для формулы (12), порядком тех приведённых разностей, которые оказываются постоянными или, говоря точнее, которые мы условно принимаем постоянными.

Рассмотрим пример интерполирования при неравноотстоящих значениях аргумента, воспользовавшись исследованием той химической реакции, которая была рассмотрена в примере 2 (стр. 280). При другой серии опытов с теми же веществами определялись значения A без строгого соблюдения постоянства интервалов t . Значения t в минутах и A в процентах, полученные при опытах, приведены в таблице XXVII.

Т а б л и ц а XXVII

t	8	11	17	27	31
A	81,2	74,8	63,0	47,0	42,6

Определим попрежнему процентное содержание вещества в системе по истечении 25 минут после начала реакции.

Составляем таблицу приведённых разностей. Результаты этих вычислений даны в таблице XXVIII, причём в ней указаны те арифметические действия, которые приходится выполнять для определения $\bar{\Delta}y$ различных порядков.

Примем условно, что третьи приведённые разности можно считать постоянными и что их среднее значение равно 0,00045. Таким образом, в формуле (14) мы можем

Таблица XXVIII

	A	$\overline{\Delta A}$	$\overline{\Delta^2 A}$	$\overline{\Delta^3 A}$
8	81,2			
11	74,8	$-6,4 : 3 = -2,13$	$0,16 : 9 = 0,018$	
17	63,0	$-11,8 : 6 = -1,97$	$0,37 : 16 = 0,023$	$0,005 : 19 = 0,00026$
27	47,0	$-16,0 : 10 = -1,60$	$0,50 : 14 = 0,036$	$0,013 : 20 = 0,00065$
31	42,6	$-4,4 : 4 = -1,10$		

ограничиться четырьмя первыми членами. Вследствие этого, переходя к числовым данным и полагая t равным 25 минутам, находим:

$$A = 81,2 + (25 - 8) \cdot (-2,13) + (25 - 8) \cdot (25 - 11) \cdot 0,018 + \\ + (25 - 8) \cdot (25 - 11) \cdot (25 - 17) \cdot 0,00045$$

или

$$A = 81,2 - 17 \cdot 2,13 + 17 \cdot 14 \cdot 0,018 + 17 \cdot 14 \cdot 8 \cdot 0,00045.$$

Производя вычисления, получаем:

$$A = 50,13\%,$$

т. е. результат, очень близкий к полученному в примере 2.

Результаты исследования данной реакции, приведённые в этих двух примерах, очень интересно сравнить графически. Если данные таблиц XXIII и XXVII изобразить графически, откладывая по оси абсцисс время t , а по оси ординат — процентное содержание A вещества в системе, то, как мы видим (рис. 47), точки первой и второй серий опытов, отмеченные на графике соответственно крестиками и кружочками, очень хорошо ложатся на кривую. Одновременно мы видим, что по истечении двадцати пяти

минут после начала реакции в системе остаётся количество вещества A , очень близкое к 50% его начального количества, что мы и получили в обоих случаях, применяя интерполирование.

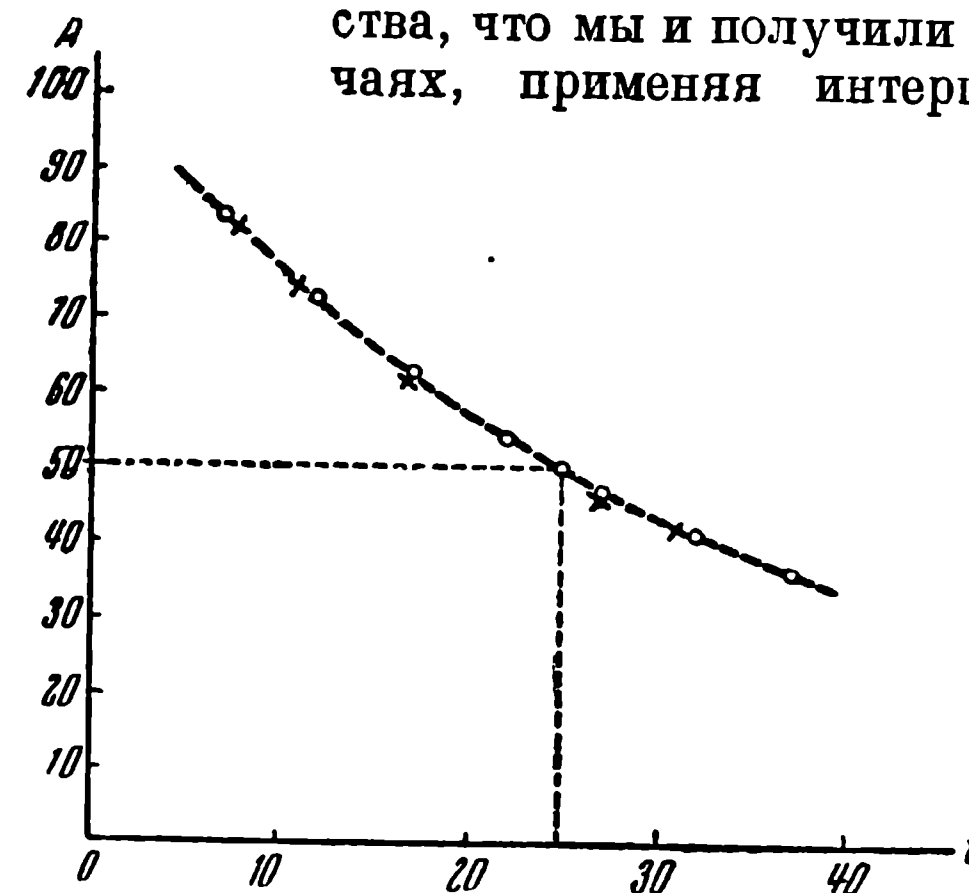


Рис. 47.

5. Точность интерполяционных формул. Кроме наиболее простых интерполяционных формул, указанных выше, было предложено ещё несколько различных интерполяционных формул, более сложных, которые часто требуют и более сложных вычислений, но вместе с тем иногда представляют несомненные преимущества. Вычисления во всех случаях начинаются с составления таблицы разностей, которая нам даёт возможность установить степень и число членов интерполирующего многочлена. Его остальные члены мы отбрасываем, и благодаря этому интерполирующая функция становится пригодной для вычислений. Таким образом, функцию $f(x)$ мы заменяем приближённой функцией $\varphi(x)$, которую всегда считаем целым рациональным многочленом с ограниченным числом членов. Геометрически такая функция изображается параболической кривой, порядок которой определяется степенью интерполирующего многочлена. Мы видим, таким образом, что кривую, определяемую уравнением (1), (стр. 257), мы заменяем на некотором протяжении отрезком параболы определённого порядка.

Из того, что было только что сказано, у нас невольно могут возникнуть следующие вопросы: во-первых, всегда ли возможна такая замена функции $f(x)$ целым рациональным многочленом или всегда ли возможна замена кривой $y = f(x)$ отрезком параболы и, во-вторых, если мы делаем такую замену, то как определить величину той ошибки, которую мы при этом допускаем.

В ответ на эти вопросы можно сказать следующее:

1) Согласно общей теореме дифференциального исчисления каждую функцию $f(x)$ можно представить в виде степенного ряда, т. е. в виде целого рационального многочлена степени n , но при условии, что эта функция является непрерывной и имеет непрерывные производные до n -й производной включительно. Эти ряды выражаются вполне определёнными формулами, так что каждый член ряда может быть вычислен по начальным данным; при этом, однако, возникают вопросы о сходимости тех рядов, которые мы получаем в результате такого разложения функций. Если мы имеем достаточно веские основания предполагать, что некоторый физический процесс, отвечающий функции $f(x)$, носит непрерывный характер (стр. 204), то замену функции $f(x)$ целым рациональным многочленом можно считать достаточно обоснованной, но при этом необходимо иметь в виду, что в данном интервале значений x ряд, который мы получаем, должен быть сходящимся. Последнее условие можно считать очевидным: интерполяционные ряды имеют всегда ограниченное число членов, и абсолютная величина последовательных членов постепенно уменьшается.

2) Если мы функцию $f(x)$ разлагаем в бесконечный степенной ряд и при вычислении значений $f(x)$ ограничиваемся его первыми n членами, то ошибка, которую мы при этом делаем, определяется величиной так называемого остаточного члена этого ряда. Тем же приёмом можно определить и величину той ошибки, которую мы допускаем, заменяя функцию $f(x)$ приближённой функцией $\varphi(x)$. Действительно, если мы возьмём разность функций $f(x)$ и $\varphi(x)$, т. е. величину

$$f(x) - \varphi(x) = \xi,$$

то эта разность и будет определять величину ошибки, которую мы допускаем при интерполировании. Величину ξ обычно так и называют ошибкой, или погрешностью интерполирования. Очевидно, что ошибка интерполирования определяется алгебраической суммой всех членов в интерполяционной формуле, которые мы отбрасываем. Мы можем определить величину ξ тем же приёмом, который применяется в дифференциальном исчислении при вычислении остаточного члена бесконечного степенного ряда. Но в этом обыкновенно необходимости не встречается, так как оценить величину ошибки интерполирования можно значительно проще, хотя и менее точно, считая, что приближённое значение ξ с достаточной точностью определяется величиной первого из отбрасываемых членов интерполяционного ряда. Таким образом, величина ошибки интерполирования определяется непосредственно на основании таблицы разностей. Определяя величину первого из отбрасываемых членов ряда, следует вводить при вычислениях среднее значение соответствующей разности, величину которой определяют как среднее арифметическое абсолютных значений всех разностей этого порядка.

Определим ошибки интерполирования в некоторых из тех примеров, которые были рассмотрены выше.

1) Из таблицы XXIV мы видим, что среднее значение разности третьего порядка, которую мы условно приняли равной нулю, можно принять равным 0,12.

Величина четвёртого члена, отброшенного нами в интерполяционном ряде, определяется выражением

$$\frac{0,12}{2 \cdot 3 \cdot 5^3} (25 - 22)(25 - 27)(25 - 32),$$

т. е. оказывается приближённо равной 0,006. Таким образом, мы можем считать, что в полученном нами результате 50,25% все цифры являются верными, и если мы это число округляем до десятых долей, то только потому, что исходные данные имеют ту же точность.

2) При вычислении плотности воды мы пользовались данными таблицы XVIII, полагая, что третьи разности

равны нулю. Считая среднее значение третьих разностей равным 0,0000003, находим величину первого из отброшенных членов ряда:

$$\frac{0,7 \cdot (0,7 - 1) \cdot (0,7 - 2)}{2 \cdot 3} \cdot 0,0000003.$$

Величина этого члена оказывается равной приближённо 0,00000001; таким образом, мы можем считать, что все шесть десятичных знаков в полученном нами значении плотности воды являются верными.

3) При определении значения эллиптического интеграла (таблица XXV) мы приняли шестые разности условно равными нулю. Среднее значение шестых разностей можно принять равным 0,00017. Таким образом, величина первого из отброшенных членов ряда в данном случае определяется выражением

$$\frac{0,5 (0,5 - 1) (0,5 - 2) (0,5 - 3) (0,5 - 4) (0,5 - 5)}{2 \cdot 3 \cdot 4 \cdot 5 \cdot 6} 17 \cdot 10^{-5},$$

т. е. оказывается приближённо равной 0,000004. Таким образом, и в этом случае мы можем считать, что в полученном нами значении эллиптического интеграла все пять десятичных знаков оказываются верными.

Как мы видим, во всех трёх случаях ошибка результата вычислений не превышает ошибки начальных значений y . Это оказывается общим правилом, которое наблюдается при всех интерполяционных вычислениях: если мы эти вычисления ведём с достаточной точностью, т. е. отбрасываем только те разности, которые являются очень малыми, практически равными нулю, то точность результата вычислений оказывается такой же, как точность исходных данных.

В заключение следует сказать несколько слов о возможности применять интерполяционные формулы за пределами того интервала значений x , в котором были произведены измерения. Такую операцию, т. е. вычисление значений y за пределами непосредственно исследованного отрезка кривой, принято называть *экстраполированием* или *экстраполяцией*. Этим приёмом следует пользоваться с очень большой осторожностью,

так как мы не можем знать поведение функции $f(x)$ за пределами исследованного интервала, поскольку аналитическое выражение этой функции остаётся нам неизвестным. Тем не менее, если кривая обнаруживает вполне плавный ход и нет оснований ожидать резких изменений в данном процессе за пределами исследованного участка кривой, то принято допускать экстраполяцию, но в очень узких пределах, например, если мы не выходим за пределы величины h , т. е. того постоянного интервала значений x , который определяется формулами (5). Экстраполяцию в более широких пределах, т. е. в более отдалённых точках кривой, следует считать всегда рискованной, так как мы можем получить совершенно неправильные значения y .

ГЛАВА IX

ОСНОВЫ ГАРМОНИЧЕСКОГО АНАЛИЗА

1. Периодические процессы. Если мы имеем две величины x и y , связанные функциональной зависимостью $y=f(x)$, то иногда оказывается, что изменения y происходят периодически, т. е. в то время как значения x мы изменяем в одном направлении, например увеличиваем, значения y периодически то возрастают, то убывают. Таким образом, можно написать:

$$y=f(x)=f(x+T)=f(x+2T)=\dots=f(x+nT),$$

где T обозначает некоторую постоянную величину, а n —произвольное целое число, положительное или отрицательное—безразлично. В таких случаях y называют периодической функцией x , а величину T называют периодом функции y .

Периодичность может наблюдаться как в пространстве, так и во времени. Особенно большое значение имеет периодичность во времени, так как такие периодические явления встречаются чрезвычайно часто: таковы, например, механические колебания, большинство акустических явлений, электрические колебания и т. д.

Периодичность функций очень наглядно обнаруживается при их графическом изображении. Значения той величины, которая обнаруживает периодические изменения, на графиках принято откладывать по оси ординат. Нулевое значение этой величины обыкновенно выбирают так, чтобы кривая располагалась по обе стороны оси абсцисс (рис. 48). Периодические кривые иногда могут иметь очень сложный вид, тем не менее, если кривая удовлетво-

ряет условию периодичности, то всегда можно найти такую точку её пересечения с осью абсцисс, начиная с которой все ординаты кривой в точности повторяются.

Периодические кривые получаются в результате сложения нескольких гармонических колебаний, что можно выполнить как графическими, так и аналитическими приемами. В последнем случае мы должны сложить все данные гармонические колебания, пользуясь теоремами сложения

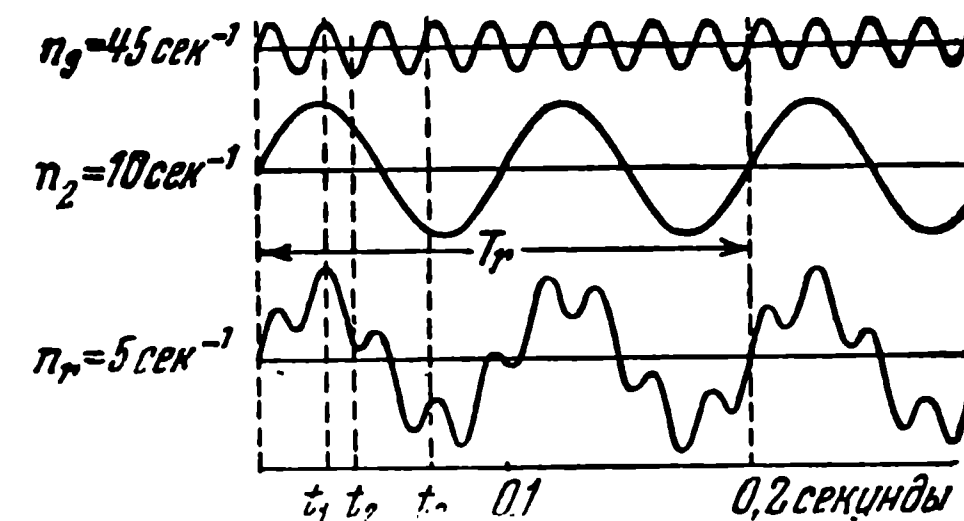


Рис. 48.

колебаний, что иногда может представлять трудность. Более простым оказывается процесс графического получения периодической кривой по данным графикам складываемых колебаний. В этом случае следует вычертить в определённом масштабе синусоиды, отвечающие всем данным гармоническим колебаниям, принимая во внимание их начальные фазы. Геометрическая сумма их ординат в каждой точке даёт, как известно, ординаты точек той периодической кривой, которая получается в результате сложения данных колебаний. Так, в нижней части рис. 48 дана периодическая кривая, которая является результатом сложения двух колебаний, изображённых в верхней части рисунка, в чём нетрудно убедиться простым суммированием их ординат в любой точке.

Обратная задача, т. е. разложение периодических кривых на ряд синусоидальных колебаний, представляется значительно более сложной. В общем случае, если число составляющих колебаний и их параметры остаются нам неизвестными, эта задача не может быть решена. Вследствие этого приходится вводить определённое ограничение,

предполагая, что составляющие колебания образуют ряд синусоидальных колебаний, из которых одно основное колебание имеет наибольший период, равный периоду функции $f(x)$. Что же касается всех остальных колебаний, то мы предполагаем, что их периоды по отношению к периоду основного колебания постепенно уменьшаются, следуя тому же закону, по которому происходит уменьшение последовательных членов так называемого гармонического ряда:

$$1 + \frac{1}{2} + \frac{1}{3} + \frac{1}{4} + \dots$$

Иначе говоря, мы предполагаем, что периоды всех составляющих колебаний пропорциональны периоду функции $f(x)$, причём коэффициентами пропорциональности служат последовательные члены ряда натуральных чисел. Таким образом, обозначая периоды составляющих колебаний через T с соответствующими индексами, мы можем написать:

$$T_0 = 2T_1 = 3T_2 = 4T_3 = \dots, \quad (1)$$

где T_0 обозначает период основного колебания, равный периоду функции $f(x)$.

Ряды колебаний, удовлетворяющие условию (1), называются гармоническими рядами колебаний. Наибольший период или наименьшая частота в таких рядах колебаний образует их основной период, или основную частоту.

Условие (1) мы вводим вследствие двух причин: во-первых, гармонические ряды колебаний, как оказывается, очень часто встречаются на практике, например, при акустических явлениях, и, во-вторых, теоретические исследования показывают, что любую периодическую кривую можно рассматривать как результат сложения бесконечно большого числа колебаний, которые образуют гармонический ряд. Последнее обстоятельство даёт возможность обосновать общее решение задачи о разложении произвольной периодической функции на составляющие, но при условии, что они образуют гармонический ряд. Теоретически эта задача допускает точное решение в общем виде,

причём следует обратить внимание на то, что определённые условия мы накладываем только на периоды гармонических составляющих, что же касается их амплитуд и начальных фаз, то здесь мы никаких ограничений не вводим.

Задачи этого рода и составляют содержание того отдела прикладной математики, который носит название гармонического анализа.

2. Гармонический анализ. В основе гармонического анализа лежит так называемая теорема Фурье. Согласно этой теореме некоторые классы функций, в том числе и функции периодические, при определённых условиях допускают разложение в виде бесконечного тригонометрического ряда, т. е. ряда, в котором независимая переменная входит под знаком тригонометрической функции, а именно, под знаком синуса и косинуса. Так как число членов ряда бесконечно велико, то при вычислении значений y по данным значениям x нам приходится отбрасывать крайние члены ряда, вводя в вычисления только некоторое число его первых членов. Таким образом, при вычислениях мы заменяем функцию $f(x)$ некоторой приближённой функцией $\varphi(x)$ совершенно аналогично тому, что имеет место при интерполировании (стр. 258). Поэтому иногда гармонический анализ называют тригонометрическим интерполированием.

При анализе периодических кривых вводят одно ограничение, которое состоит в том, что период функции предполагают равным 2π . Конечно, далеко не все периодические функции удовлетворяют этому условию, но нетрудно показать, что любую периодическую функцию с произвольным периодом можно привести к периоду, равному 2π , при помощи простой замены независимой переменной.

Действительно, допустим, что периодическая функция

$$y = f(x)$$

имеет некоторый период θ , отличный от 2π . Введём новую независимую переменную z , полагая, что

$$z = a + bx, \quad (2)$$

и найдём такие значения a и b , чтобы при изменении x от 0 до θ z изменялось от 0 до 2π .

Подставляя предельные значения x и z в последнее выражение, получаем два уравнения:

$$0 = a + b \cdot 0 \text{ и } 2\pi = a + b\theta.$$

Отсюда находим:

$$a = 0, \quad b = \frac{2\pi}{\theta}.$$

Таким образом, значение z оказывается равным:

$$z = \frac{2\pi}{\theta} x.$$

Отсюда имеем:

$$x = \frac{\theta}{2\pi} z.$$

Мы видим, что если вместо x ввести новую независимую переменную, определяемую последним выражением, то период функции становится равным 2π . Поэтому мы будем предполагать, что периодическая функция $f(x)$ имеет период, равный 2π , причём это предположение, очевидно, нисколько не ограничивает общности всех дальнейших выводов.

Допустим, что для периодической функции

$$y = f(x)$$

с периодом, равным 2π , дан ряд соответственных значений x и y :

$$\begin{aligned} x_0, x_1, x_2, x_3, \dots, x_n, \\ y_0, y_1, y_2, y_3, \dots, y_n. \end{aligned} \quad (3)$$

Значения x мы будем считать равноотстоящими, при этом предполагаем, что они охватывают весь интервал, отвечающий полному периоду функции, т. е. 2π . Иначе говоря, если мы возьмём ряд следующих значений x с тем же интервалом, равных $x_{n+1}, x_{n+2}, x_{n+3}, \dots$, то y получает свои прежние значения, равные соответственно

$$y_0, y_1, y_2, \dots$$

Согласно общей теореме Фурье функцию $f(x)$ можно представить в форме бесконечного тригонометрического ряда такого вида:

$$y = f(x) = a_0 + a_1 \cos x + a_2 \cos 2x + a_3 \cos 3x + a_4 \cos 4x + \dots \\ \dots + b_1 \sin x + b_2 \sin 2x + b_3 \sin 3x + b_4 \sin 4x + \dots \quad (4)$$

Этот ряд иногда пишут также в такой форме:

$$y = f(x) = a_0 + (a_1 \cos x + b_1 \sin x) + (a_2 \cos 2x + b_2 \sin 2x) + \\ + (a_3 \cos 3x + b_3 \sin 3x) + (a_4 \cos 4x + b_4 \sin 4x) + \dots \quad (4')$$

Нетрудно видеть, что выражения, стоящие в скобках в правой части последнего ряда, представляют собой отдельные гармонические составляющие, или гармоники, так как каждое из них может быть приведено к обычному виду уравнения гармонического колебания следующим приёмом. Обозначая выражения, стоящие в скобках, через c , в общем виде можно написать:

$$c_n = a_n \cos nx + b_n \sin nx. \quad (4'')$$

В этом выражении полагаем:

$$a_n = A_n \sin \varphi_n \text{ и } b_n = A_n \cos \varphi_n. \quad (5)$$

Такая подстановка, очевидно, всегда возможна. Действительно, если последние уравнения сначала возвести во вторую степень и сложить, а затем разделить одно на другое, то получаем такие выражения:

$$a_n^2 + b_n^2 = A_n^2 \text{ и } \frac{a_n}{b_n} = \operatorname{tg} \varphi_n. \quad (5')$$

Отсюда видно, что, каковы бы ни были значения a_n и b_n , мы всегда сумеем найти такие величины A_n и φ_n , что последние равенства будут удовлетворены и указанная подстановка окажется возможной.

Подставляя значения a_n и b_n в выражение (4''), находим:

$$c_n = A_n (\sin \varphi_n \cos nx + \cos \varphi_n \sin nx)$$

или

$$c_n = A_n \sin (nx + \varphi_n),$$

что вполне отвечает обычному виду уравнения гармонического колебания:

$$s = A \sin \left(2\pi \frac{t}{T} + \varphi_0 \right),$$

причём A_n и φ_n обозначают амплитуду и начальную фазу в этом колебании, а nx оказывается равным:

$$nx = 2\pi \frac{t}{T}.$$

Полагая в этом выражении n равным последовательно 1, 2, 3, 4, ... и обозначая периоды этих колебаний T индексами 0, 1, 2, 3 и т. д., находим:

$$x = 2\pi \frac{t}{T_0}; \quad 2x = 2\pi \frac{t}{T_1}; \quad 3x = 2\pi \frac{t}{T_2}; \quad 4x = 2\pi \frac{t}{T_3}; \dots$$

Отсюда имеем:

$$T_0 = 2T_1 = 3T_2 = 4T_3 = \dots$$

Мы видим, что все члены ряда (4) удовлетворяют условию (1).

Для дальнейшего следует обратить внимание на то, что A_n и φ_n в выражениях (5) оказываются амплитудой и начальной фазой отдельных колебаний и что их величина определяется формулами (5'), в которые входят соответствующие коэффициенты a и b ряда (4).

В зависимости от степени приближения мы можем брать различное число членов ряда (4), отбрасывая все остальные. Таким образом, допуская, что мы ограничиваемся первыми $(2n+1)$ членами этого ряда, и вводя вместо функции $f(x)$ приближённую функцию $\varphi(x)$, мы можем написать:

$$\begin{aligned} y \approx \varphi(x) = & a_0 + a_1 \cos x + a_2 \cos 2x + a_3 \cos 3x + \dots \\ & \dots + a_n \cos nx + b_1 \sin x + b_2 \sin 2x + b_3 \sin 3x + \dots \\ & \dots + b_n \sin nx. \end{aligned} \quad (6).$$

Ошибка ξ , которую мы при этом делаем, определяется величиной разности:

$$\xi = f(x) - \varphi(x). \quad (7)$$

Наша задача состоит в том, чтобы, во-первых, определить число членов ряда (6), которым можно ограничиться, и, во-вторых, найти такие значения его коэффициентов $a_0, a_1, a_2, \dots, a_n, b_1, b_2, b_3, \dots, b_n$, при которых ошибка, определяемая выражением (7), имеет наименьшую величину. Общий метод решения этих вопросов можно изложить следующим образом.

Подставляя в выражение (7) последовательно все значения x , мы определим для каждого из них в отдельности ту ошибку, которая при этом получается, т. е. определим величину разности (7) для каждого из значений x . Сумма квадратов всех этих ошибок должна иметь наименьшее значение (стр. 358). Это условие может служить для определения тех значений коэффициентов уравнения (6), при которых ошибка, определяемая выражением (7), имеет наименьшее значение. Поэтому, выражая сумму квадратов всех ошибок в виде определённого интеграла с пределами 0 и 2π , можем написать:

$$\int_0^{2\pi} [f(x) - \varphi(x)]^2 dx = \text{минимум}.$$

Для того чтобы определить условия минимума этого выражения, следует найти его частные производные по всем коэффициентам, которые входят в подинтегральное выражение, и приравнять эти производные нулю. В подинтегральное выражение входят все коэффициенты многочлена (6), число которых равно $(2n+1)$. Таким образом, выполняя дифференцирование и приравнивая производные нулю, мы получим $(2n+1)$ уравнений, которые могут служить для определения всех коэффициентов многочлена (6). Эти уравнения в общем виде определяются такими выражениями:

$$\begin{aligned} \int_0^{2\pi} [f(x) - \varphi(x)] \sin mx dx &= 0, \\ \int_0^{2\pi} [f(x) - \varphi(x)] \cos mx dx &= 0, \end{aligned}$$

где m имеет значения всех целых чисел от 0 до n .

Эти уравнения допускают ряд упрощений, в результате которых для коэффициентов a и b получаются такие общие выражения:

1) Для свободного члена a_0 :

$$a_0 = \frac{1}{2\pi} \int_0^{2\pi} f(x) dx.$$

2) Для остальных коэффициентов a :

$$a_n = \frac{1}{\pi} \int_0^{2\pi} f(x) \cos nx dx.$$

3) Для коэффициентов b :

$$b_n = \frac{1}{\pi} \int_0^{2\pi} f(x) \sin nx dx,$$

где n имеет значения всех целых чисел от 1 до n .

Определение коэффициентов a и b этим методом на практике не всегда оказывается возможным, так как интегралы в правой части этих уравнений не всегда поддаются точному вычислению, кроме того, алгебраическое выражение функции $f(x)$ часто остаётся нам неизвестным. Вследствие этого очень часто применяют некоторые практические приёмы, которые оказываются значительно более простыми.

Условия, при которых вычисления коэффициентов ряда (6) выполняются без больших затруднений, состоят в следующем:

1) Период функции $f(x)$, равный 2π , делят на чётное число равных частей и находят значения y для каждой точки деления. Если дан график периодической кривой, то выполнить это условие нетрудно, следует просто измерить ординаты, построенные для точек деления. Чётное число точек деления, или, как обычно говорят, чётное число ординат, следует брать потому, что при этом условии значения синуса и косинуса по абсолютной величине повторяются в каждом квадранте, что очень об-

легчает вычисления. Ещё удобнее брать число ординат, кратное четырём.

2) Общее число членов многочлена (6) берут равным числу ординат. Так как в этом многочлене первый член свободен от x , то при чётном числе ординат число членов, содержащих синус, и число членов, содержащих косинус, должны отличаться на единицу. Обыкновенно последний член, содержащий синус, отбрасывают. Таким образом, если мы взяли $2n$ ординат и соответственно этому в многочлене (6) $2n$ членов, то его следует написать, отбрасывая последний член, в таком виде:

$$\begin{aligned} y \approx \varphi(x) = & a_0 + a_1 \cos x + a_2 \cos 2x + a_3 \cos 3x + \dots \\ & \dots + a_n \cos nx + b_1 \sin x + b_2 \sin 2x + b_3 \sin 3x + \dots \\ & \dots + b_{n-1} \sin (n-1)x. \quad (6') \end{aligned}$$

При анализе различных периодических кривых на практике берут различное число ординат, например 6, 8, 12 и т. д. до 24 ординат; последнее число является обычно предельным для непосредственных вычислений. Наиболее часто применяются при вычислениях двенадцать ординат; вследствие этого мы выведем формулы и рассмотрим схему вычислений только для этого случая.

3. Схема вычислений для двенадцати ординат. При 12 ординатах мы можем в многочлене (6') взять, кроме свободного члена, шесть членов, содержащих косинус, и пять членов, содержащих синус. Таким образом, для случая 12 ординат многочлен (6') получает вид

$$\begin{aligned} y \approx \varphi(x) = & a_0 + a_1 \cos x + a_2 \cos 2x + a_3 \cos 3x + a_4 \cos 4x + \\ & + a_5 \cos 5x + a_6 \cos 6x + b_1 \sin x + b_2 \sin 2x + b_3 \sin 3x + \\ & + b_4 \sin 4x + b_5 \sin 5x. \quad (6'') \end{aligned}$$

Так как период функции y , равный 2π , разделён на 12 частей, то интервал для значений x равняется 30° , т. е. x получает значения $0^\circ, 30^\circ, 60^\circ, 90^\circ, \dots, 330^\circ$. При значениях x , равных $360^\circ, 390^\circ, 420^\circ, \dots$, значения y начинают повторяться, т. е. y получает значения, равные соответственно $y_0, y_1, y_2, \dots$. Таким образом, мы

получаем следующую таблицу соответственных значений x и y :

Таблица XXIX

x	0°	30°	60°	90°	120°	150°	180°	210°	240°	270°	300°	330°	360°
y	y_0	y_1	y_2	y_3	y_4	y_5	y_6	y_7	y_8	y_9	y_{10}	y_{11}	y_0

Все двенадцать соответственных значений x и y подставляем последовательно в уравнение (6*). В результате получаем 12 уравнений, в которых числовыми коэффициентами являются значения синуса и косинуса от аргументов, или равных 0° , или кратных 30° . Во всех случаях числовые значения коэффициентов оказываются равными:

$$0, \pm 1, \pm \frac{1}{2} \text{ и } \pm \frac{\sqrt{3}}{2},$$

т. е. выражаются одними и теми же числами. Отсюда становится понятным, почему рекомендуется брать число ординат, кратное четырём: при этом условии значения некоторых аргументов оказываются равными $0^\circ, 90^\circ, 180^\circ, 270^\circ, 360^\circ$ и т. д., и вычисления сильно упрощаются. В данном случае получаются такие уравнения:

$$y_0 = a_0 + a_1 + a_2 + a_3 + a_4 + a_5 + a_6,$$

$$y_1 = a_0 + \frac{\sqrt{3}}{2} a_1 + \frac{1}{2} a_2 - \frac{1}{2} a_4 - \frac{\sqrt{3}}{2} a_5 - a_6 + \\ + \frac{1}{2} b_1 + \frac{\sqrt{3}}{2} b_2 + b_3 + \frac{\sqrt{3}}{2} b_4 + \frac{1}{2} b_5,$$

$$y_2 = a_0 + \frac{1}{2} a_1 - \frac{1}{2} a_2 - a_3 - \frac{1}{2} a_4 + \frac{1}{2} a_5 + a_6 + \\ + \frac{\sqrt{3}}{2} b_1 + \frac{\sqrt{3}}{2} b_2 - \frac{\sqrt{3}}{2} b_4 - \frac{\sqrt{3}}{2} b_5,$$

$$y_3 = a_0 - a_2 + a_4 - a_6 + b_1 - b_3 + b_5,$$

$$y_4 = a_0 - \frac{1}{2} a_1 - \frac{1}{2} a_2 + a_3 - \frac{1}{2} a_4 - \frac{1}{2} a_5 + a_6 + \\ + \frac{\sqrt{3}}{2} b_1 - \frac{\sqrt{3}}{2} b_2 + \frac{\sqrt{3}}{2} b_4 - \frac{\sqrt{3}}{2} b_5,$$

$$y_5 = a_0 - \frac{\sqrt{3}}{2} a_1 + \frac{1}{2} a_2 - \frac{1}{2} a_4 + \frac{\sqrt{3}}{2} a_5 - a_6 + \\ + \frac{1}{2} b_1 - \frac{\sqrt{3}}{2} b_2 + b_3 - \frac{\sqrt{3}}{2} b_4 - \frac{1}{2} b_5,$$

$$y_6 = a_0 - a_1 + a_2 - a_3 + a_4 - a_5 + a_6,$$

$$y_7 = a_0 - \frac{\sqrt{3}}{2} a_1 + \frac{1}{2} a_2 - \frac{1}{2} a_4 + \frac{\sqrt{3}}{2} a_5 - a_6 - \\ - \frac{1}{2} b_1 + \frac{\sqrt{3}}{2} b_2 - b_3 + \frac{\sqrt{3}}{2} b_4 - \frac{1}{2} b_5,$$

$$y_8 = a_0 - \frac{1}{2} a_1 - \frac{1}{2} a_2 + a_3 - \frac{1}{2} a_4 - \frac{1}{2} a_5 + a_6 - \\ - \frac{\sqrt{3}}{2} b_1 + \frac{\sqrt{3}}{2} b_2 - \frac{\sqrt{3}}{2} b_4 + \frac{\sqrt{3}}{2} b_5,$$

$$y_9 = a_0 - a_2 + a_4 - a_6 - b_1 + b_3 - b_5,$$

$$y_{10} = a_0 + \frac{1}{2} a_1 - \frac{1}{2} a_2 - a_3 - \frac{1}{2} a_4 + \frac{1}{2} a_5 + a_6 - \\ - \frac{\sqrt{3}}{2} b_1 - \frac{\sqrt{3}}{2} b_2 + \frac{\sqrt{3}}{2} b_4 + \frac{\sqrt{3}}{2} b_5,$$

$$y_{11} = a_0 + \frac{\sqrt{3}}{2} a_1 + \frac{1}{2} a_2 - \frac{1}{2} a_4 - \frac{\sqrt{3}}{2} a_5 - a_6 - \\ - \frac{1}{2} b_1 - \frac{\sqrt{3}}{2} b_2 - b_3 - \frac{\sqrt{3}}{2} b_4 - \frac{1}{2} b_5.$$

Для того чтобы из этих уравнений определить значения коэффициентов $a_0, a_1, a_2, \dots, a_6$ и $b_1, b_2, \dots, b_5$, поступаем следующим образом:

1) Для определения a_0 складываем все уравнения. Все члены в правой части полученного уравнения, за исключением членов, содержащих a_0 , взаимно уничтожаются, и мы получаем:

$$12a_0 = y_0 + y_1 + y_2 + y_3 + y_4 + y_5 + y_6 + y_7 + \\ + y_8 + y_9 + y_{10} + y_{11}.$$

2) Для определения a_1 каждое из уравнений умножаем на числовой коэффициент, который стоит при a_1 в этом уравнении, и все полученные уравнения

складываем. После приведения получаем:

$$6a_1 = y_0 + \frac{\sqrt{3}}{2} y_1 + \frac{1}{2} y_2 - \frac{1}{2} y_4 - \frac{\sqrt{3}}{2} y_5 - y_6 - \\ - \frac{\sqrt{3}}{2} y_7 - \frac{1}{2} y_8 + \frac{1}{2} y_{10} + \frac{\sqrt{3}}{2} y_{11}.$$

Точно так же поступаем для определения всех остальных коэффициентов, т. е. для определения какого-либо коэффициента умножаем каждое из уравнений на числовой коэффициент, который стоит в этом уравнении при определяемом коэффициенте, затем все полученные уравнения складываем и делаем приведение. В результате получаются следующие выражения:

$$6a_2 = y_0 + \frac{1}{2} y_1 + \frac{1}{2} y_2 - y_3 - \frac{1}{2} y_4 + \frac{1}{2} y_5 + y_6 + \\ + \frac{1}{2} y_7 - \frac{1}{2} y_8 - y_9 - \frac{1}{2} y_{10} + \frac{1}{2} y_{11},$$

$$6a_3 = y_0 - y_2 + y_4 - y_6 + y_8 - y_{10},$$

$$6a_4 = y_0 - \frac{1}{2} y_1 - \frac{1}{2} y_2 + y_3 - \frac{1}{2} y_4 - \frac{1}{2} y_5 + y_6 - \\ - \frac{1}{2} y_7 - \frac{1}{2} y_8 + y_9 - \frac{1}{2} y_{10} - \frac{1}{2} y_{11},$$

$$6a_5 = y_0 - \frac{\sqrt{3}}{2} y_1 + \frac{1}{2} y_2 - \frac{1}{2} y_4 + \frac{\sqrt{3}}{2} y_5 - y_6 + \\ + \frac{\sqrt{3}}{2} y_7 - \frac{1}{2} y_8 + \frac{1}{2} y_{10} - \frac{\sqrt{3}}{2} y_{11},$$

$$12a_6 = y_0 - y_1 + y_2 - y_3 + y_4 - y_5 + \\ + y_6 - y_7 + y_8 - y_9 + y_{10} - y_{11},$$

$$6b_1 = \frac{1}{2} y_1 + \frac{\sqrt{3}}{2} y_2 + y_3 + \frac{\sqrt{3}}{2} y_4 + \frac{1}{2} y_5 - \frac{1}{2} y_7 - \\ - \frac{\sqrt{3}}{2} y_8 - y_9 - \frac{\sqrt{3}}{2} y_{10} - \frac{1}{2} y_{11},$$

$$6b_2 = \frac{\sqrt{3}}{2} (y_1 + y_2 - y_4 - y_5 + y_7 + y_8 - y_{10} - y_{11}),$$

$$6b_3 = y_1 - y_3 + y_5 - y_7 + y_9 - y_{11},$$

$$6b_4 = \frac{\sqrt{3}}{2} (y_1 - y_2 + y_4 - y_5 + y_7 - y_8 + y_{10} - y_{11}),$$

$$6b_5 = \frac{1}{2} y_1 - \frac{\sqrt{3}}{2} y_2 + y_3 - \frac{\sqrt{3}}{2} y_4 + \frac{1}{2} y_5 - \\ - \frac{1}{2} y_7 + \frac{\sqrt{3}}{2} y_8 - y_9 + \frac{\sqrt{3}}{2} y_{10} - \frac{1}{2} y_{11}.$$

Из полученных выражений можно непосредственно вычислить значения всех коэффициентов многочлена (6*), но это вычисление является утомительным и его заменяют более простыми вычислениями, пользуясь тем, что в правые части этих равенств входят суммы и разности одних и тех же парных значений y . Действительно, при помощи группировки членов в правых частях всех последних равенств приводим их к следующему виду:

$$12a_0 = (y_0 + y_6) + (y_1 + y_{11}) + (y_2 + y_{10}) + (y_3 + y_9) + \\ + (y_4 + y_8) + (y_5 + y_7),$$

$$6a_1 = (y_0 - y_6) + \frac{\sqrt{3}}{2} (y_1 + y_{11}) + \frac{1}{2} (y_2 + y_{10}) - \\ - \frac{1}{2} (y_4 + y_8) - \frac{\sqrt{3}}{2} (y_5 + y_7),$$

$$6a_2 = (y_0 + y_6) + \frac{1}{2} (y_1 + y_{11}) - \frac{1}{2} (y_2 + y_{10}) - \\ - (y_3 + y_9) - \frac{1}{2} (y_4 + y_8) + \frac{1}{2} (y_5 + y_7),$$

$$6a_3 = (y_0 - y_6) - (y_2 + y_{10}) + (y_4 + y_8),$$

$$6a_4 = (y_0 + y_6) - \frac{1}{2} (y_1 + y_{11}) - \frac{1}{2} (y_2 + y_{10}) + \\ + (y_3 + y_9) - \frac{1}{2} (y_4 + y_8) - \frac{1}{2} (y_5 + y_7),$$

$$6a_5 = (y_0 + y_6) - \frac{\sqrt{3}}{2} (y_1 + y_{11}) + \frac{1}{2} (y_2 + y_{10}) - \\ - \frac{1}{2} (y_4 + y_8) + \frac{\sqrt{3}}{2} (y_5 + y_7),$$

$$12a_6 = (y_0 + y_6) - (y_1 + y_{11}) + (y_2 + y_{10}) - \\ - (y_3 + y_9) + (y_4 + y_8) - (y_5 + y_7),$$

$$\begin{aligned}
6b_1 &= \frac{1}{2}(y_1 - y_{11}) + \frac{\sqrt{3}}{2}(y_2 - y_{10}) + (y_3 - y_9) + \\
&\quad + \frac{\sqrt{3}}{2}(y_4 - y_8) + \frac{1}{2}(y_5 - y_7), \\
6b_2 &= \frac{\sqrt{3}}{2}[(y_1 - y_{11}) + (y_2 - y_{10}) - (y_4 - y_8) - (y_5 - y_7)], \\
6b_3 &= (y_1 - y_{11}) - (y_3 - y_9) + (y_5 - y_7), \\
6b_4 &= \frac{\sqrt{3}}{2}[(y_1 - y_{11}) - (y_2 - y_{10}) + (y_4 - y_8) - (y_5 - y_7)], \\
6b_5 &= \frac{1}{2}(y_1 - y_{11}) - \frac{\sqrt{3}}{2}(y_2 - y_{10}) + (y_3 - y_9) - \\
&\quad - \frac{\sqrt{3}}{2}(y_4 - y_8) + \frac{1}{2}(y_5 - y_7).
\end{aligned}$$

Введём новые обозначения для сумм и разностей значений y , которые повторяются в правых частях этих равенств, полагая:

$$\begin{aligned}
y_0 + y_6 &= u_0, & y_0 - y_6 &= v_0, \\
y_1 + y_{11} &= u_1, & y_1 - y_{11} &= v_1, \\
y_2 + y_{10} &= u_2, & y_2 - y_{10} &= v_2, \\
y_3 + y_9 &= u_3, & y_3 - y_9 &= v_3, \\
y_4 + y_8 &= u_4, & y_4 - y_8 &= v_4, \\
y_5 + y_7 &= u_5, & y_5 - y_7 &= v_5.
\end{aligned}$$

Подставляя в последние равенства эти обозначения, приводим их к более простому виду, причём оказывается, что значения u и v в равенствах вновь входят в виде повторяющихся парных сумм и разностей. Поэтому новой группировкой членов в правых частях равенств приводим их к такому виду:

$$\begin{aligned}
12a_0 &= (u_0 + u_3) + (u_1 + u_5) + (u_2 + u_4), \\
6a_1 &= v_0 + \frac{\sqrt{3}}{2}(u_1 - u_5) + \frac{1}{2}(u_2 - u_4), \\
6a_2 &= (u_0 - u_3) + \frac{1}{2}(u_1 + u_5) - \frac{1}{2}(u_2 + u_4), \\
6a_3 &= v_0 - (u_2 - u_4),
\end{aligned}$$

$$\begin{aligned}
6a_4 &= (u_0 + u_3) - \frac{1}{2}(u_1 + u_5) - \frac{1}{2}(u_2 + u_4), \\
6a_5 &= v_0 - \frac{\sqrt{3}}{2}(u_1 - u_5) + \frac{1}{2}(u_2 - u_4), \\
12a_6 &= (u_0 - u_3) - (u_1 + u_5) + (u_2 + u_4), \\
6b_1 &= \frac{1}{2}(v_1 + v_5) + \frac{\sqrt{3}}{2}(v_2 + v_4) + v_3, \\
6b_2 &= \frac{\sqrt{3}}{2}[(v_1 - v_5) + (v_2 - v_4)], \\
6b_3 &= (v_1 + v_5) - v_3, \\
6b_4 &= \frac{\sqrt{3}}{2}[(v_1 - v_5) - (v_2 - v_4)], \\
6b_5 &= \frac{1}{2}(v_1 + v_5) - \frac{\sqrt{3}}{2}(v_2 + v_4) + v_3.
\end{aligned}$$

Вводим опять новые обозначения, полагая:

$$\begin{aligned}
u_0 + u_3 &= r_0; & u_0 - u_3 &= s_0; \\
u_1 + u_4 &= r_1; & u_1 - u_4 &= s_1; \\
u_2 + u_5 &= r_2; & u_2 - u_5 &= s_2; \\
v_1 + v_5 &= p_1; & v_1 - v_5 &= q_1; \\
v_2 + v_4 &= p_2; & v_2 - v_4 &= q_2.
\end{aligned}$$

Вводя эти обозначения в последние равенства, приводим их к ещё более простому виду, причём оказывается, что в этих равенствах вновь повторяются одни и те же суммы и разности величин r и q . Поэтому, группируя соответствующие члены в правых частях равенств, получаем:

$$\begin{aligned}
12a_0 &= r_0 + (r_1 + r_2), \\
6a_1 &= v_0 + \frac{\sqrt{3}}{2}s_1 + \frac{1}{2}s_2, & 6b_1 &= v_3 + \frac{1}{2}p_1 + \frac{\sqrt{3}}{2}p_2, \\
6a_2 &= s_0 + \frac{1}{2}(r_1 - r_2), & 6b_2 &= \frac{\sqrt{3}}{2}(q_1 + q_2), \\
6a_3 &= v_0 - s_2, & 6b_3 &= p_1 - v_3, \\
6a_4 &= r_0 - \frac{1}{2}(r_1 + r_2), & 6b_4 &= \frac{\sqrt{3}}{2}(q_1 - q_2), \\
6a_5 &= v_0 - \frac{\sqrt{3}}{2}s_1 + \frac{1}{2}s_2, & 6b_5 &= v_3 + \frac{1}{2}p_1 - \frac{\sqrt{3}}{2}p_2, \\
12a_6 &= s_0 - (r_1 - r_2).
\end{aligned}$$

Вводим ещё раз новые обозначения, полагая:

$$\begin{aligned} r_1 + r_2 &= l, & q_1 + q_2 &= d, \\ r_1 - r_2 &= m, & q_1 - q_2 &= h. \end{aligned}$$

Подставляя эти обозначения в последние равенства и определяя в них значения коэффициентов $a_0, a_1, \dots, a_6$ и $b_1, b_2, \dots, b_5$, находим такие выражения:

$$\left. \begin{aligned} a_0 &= \frac{1}{12}(r_0 + l), \\ a_1 &= \frac{1}{6}\left(v_0 + \frac{\sqrt{3}}{2}s_1 + \frac{1}{2}s_2\right), \\ a_2 &= \frac{1}{6}\left(s_0 + \frac{1}{2}m\right), \\ a_3 &= \frac{1}{6}(v_0 - s_2), \\ a_4 &= \frac{1}{6}\left(r_0 - \frac{1}{2}l\right), \\ a_5 &= \frac{1}{6}\left(v_0 - \frac{\sqrt{3}}{2}s_1 + \frac{1}{2}s_2\right), \\ a_6 &= \frac{1}{12}(s_0 - m), \\ b_1 &= \frac{1}{6}\left(v_3 + \frac{1}{2}p_1 + \frac{\sqrt{3}}{3}p_2\right), \\ b_2 &= \frac{\sqrt{3}}{12}d, \\ b_3 &= \frac{1}{6}(p_1 - v_3), \\ b_4 &= \frac{\sqrt{3}}{12}h, \\ b_5 &= \frac{1}{6}\left(v_3 + \frac{1}{2}p_1 - \frac{\sqrt{3}}{2}p_2\right). \end{aligned} \right\} \quad (8)$$

Эти формулы и применяют для определения коэффициентов многочлена (6*), вычисляя предварительно значение величин, которые входят в правые части выражений (8).

Вычисление этих величин удобно производить, пользуясь следующей схемой:

	y_0	y_1	y_2	y_3	y_4	y_5
	y_6	y_{11}	y_{10}	y_9	y_8	y_7
Суммы (u)	u_0	u_1	u_2	u_3	u_4	u_5
Разности (v)	v_0	v_1	v_2	v_3	v_4	v_5
	u_0	u_1	u_2		v_1	v_2
	u_3	u_5	u_4		v_5	v_4
Суммы (r, p)	r_0	r_1	r_2		p_1	p_2
Разности (s, q)	s_0	s_1	s_2		q_1	q_2
		r_1	q_1			
		r_2	q_2			
Суммы (l, d)		l	d			
Разности (m, h)		m	h			

В этой схеме те величины, которые входят в окончательные формулы (8), напечатаны жирным шрифтом.

При вычислении коэффициентов a и b приходится выполнять много арифметических действий, так что всегда возможны ошибки и опiski. Поэтому принято проверять вычисления при помощи контрольных формул. Обычно применяют две контрольные формулы. Одной из них служит первое начальное уравнение.

Вторую контрольную формулу получают, вычитая двенадцатое начальное уравнение из второго. Таким образом, первая контрольная формула служит для проверки значений a , вторая для проверки значений b . Эти формулы имеют такой вид:

$$y_0 = a_0 + a_1 + a_2 + a_3 + a_4 + a_5 + a_6,$$

$$y_1 - y_{11} = v_1 = b_1 + b_5 + 2b_3 + \sqrt{3}(b_2 + b_4).$$

Вычисляя коэффициенты a и b для случая 12 ординат, мы находим в многочлене (6'), кроме свободного члена a_0 , ещё шесть следующих коэффициентов a и пять коэффициентов b . Это даёт нам пять первых гармонических составляющих, для которых мы можем по формулам (5') определить амплитуды и начальные фазы. Что же касается шестой гармонической составляющей, то для неё мы имеем только член, содержащий косинус, так как величина b_6

остаётся неизвестной; вследствие этого амплитуду и начальную фазу этой гармонической составляющей мы определить не можем. Таким образом, если нам необходимо определить шестую и следующие гармонические составляющие, то приходится брать соответственно большее число ординат. В этих случаях определение коэффициентов производится по тому же методу и отличается только иным расположением вычислений в соответствии с иными значениями числовых коэффициентов в начальных уравнениях.

Следует также заметить, что приведённая выше схема вычислений является не единственной, и было предложено ещё несколько других схем, которые, однако, нет оснований считать более простыми или более удобными.

Рассмотрим два примера гармонического анализа периодических функций для случая 12 ординат.

4. Примеры вычисления коэффициентов разложения при двенадцати ординатах.

Пример 1. В таблице XXX даны значения периодической функции y , отвечающие ряду значений аргумента x

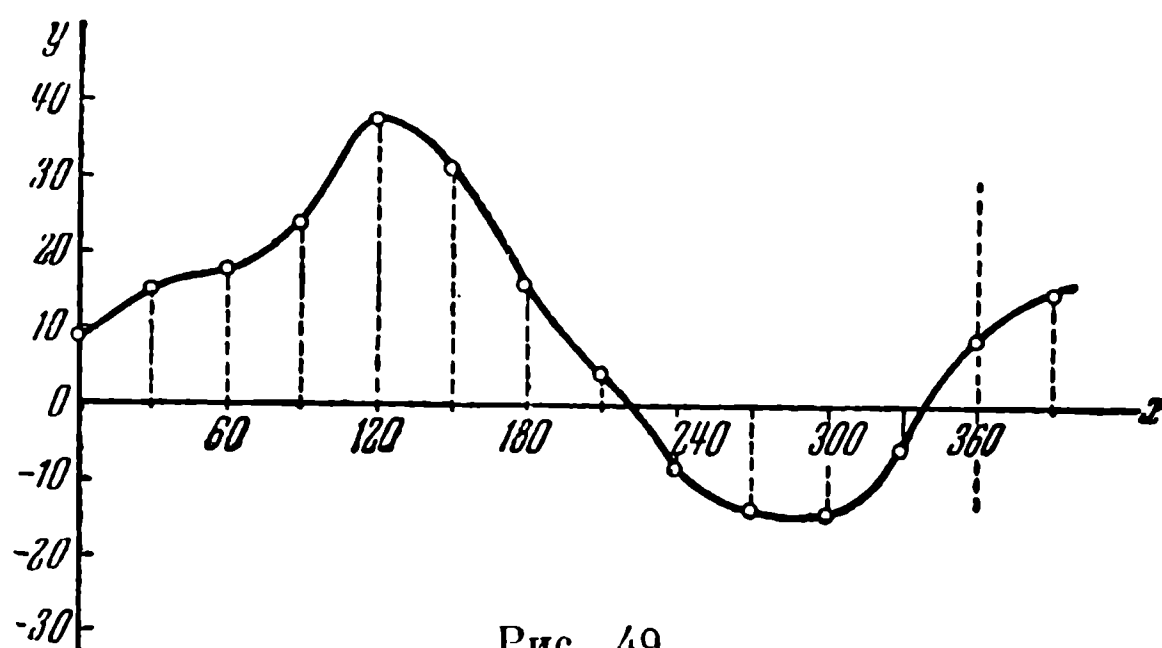


Рис. 49.

с интервалом 30° . График соответствующей кривой дан на рис. 49.

Определим коэффициенты многочлена (6*) для этого случая. Значения y располагаем по приведённой выше

Таблица XXX

x	0°	30°	60°	90°	120°	150°	180°	210°	240°	270°	300°	330°
y	9,3	15,0	17,4	23,0	37,0	31,0	15,3	4,0	-8,0	-13,2	-14,2	-6,0

схеме и находим их суммы u и разности v :

	(y)	9,3	15,0	17,4	23,0	37,0	31,0
		15,3	-6,0	-14,2	-13,2	-8,0	4,0
Суммы (u)		24,6	9,0	3,2	9,8	29,0	35,0
Разности (v)		-6,0	21,0	31,6	36,2	45,0	27,0

Располагаем по схеме значения u и v и находим их суммы r и p и их разности s и q :

(u)	24,6	9,0	3,2
	9,8	35,0	29,0
Суммы (r)	34,4	44,0	32,2
Разности (s)	14,8	— 26,0	— 25,8
(v)	21,0	31,6	
	27,0	45,0	
Суммы (p)	48,0	76,6	
Разности (q)	— 6,0	— 13,4	

Располагаем по схеме r и q и находим их суммы l и d и разности m и h :

(r)	44,0	(q)	-6,0
	32,2		-13,4
(l)	72,2	(d)	-19,4
(m)	11,8	(h)	7,4

Подставляя эти величины в уравнения (8), находим значения всех коэффициентов многочлена (6*):

$$a_0 = \frac{1}{12} (34,4 + 76,2) = 9,22,$$

$$a_1 = \frac{1}{6} \left(-6,0 - 26 \frac{\sqrt{3}}{2} - 12,9 \right) = -6,90,$$

$$a_2 = \frac{1}{6} (14,8 + 5,9) = 3,45,$$

$$a_3 = \frac{1}{6} (-6,0 + 25,8) = 3,30,$$

$$a_4 = \frac{1}{6} (34,4 - 38,1) = -0,62,$$

$$a_5 = \frac{1}{6} \left(-6,0 + 26,0 \cdot \frac{\sqrt{3}}{2} - 12,9 \right) = 0,60,$$

$$a_6 = \frac{1}{12} (14,8 - 11,8) = 0,25,$$

$$b_1 = \frac{1}{6} (36,2 + 24,0 + 66,3) = 21,09,$$

$$b_2 = \frac{\sqrt{3}}{12} (-19,4) = -2,80,$$

$$b_3 = \frac{1}{6} (48,0 - 36,2) = 1,97,$$

$$b_4 = \frac{\sqrt{3}}{12} \cdot 7,4 = 1,07,$$

$$b_5 = \frac{1}{6} (36,2 + 24,0 - 66,3) = -1,02.$$

Применяя контрольные формулы, имеем:

$$y_0 = 9,3 = 9,22 - 6,90 + 3,45 + 3,30 - 0,62 + \\ + 0,60 + 0,25 = 9,30,$$

$$y_1 - y_{11} = v_1 = 21,0 = 21,09 - 1,02 + 2 \cdot 1,97 + \\ + \sqrt{3} \cdot (-2,80 + 1,07) = 21,01.$$

Таким образом, мы видим, что значения коэффициентов a и b вычислены правильно и многочлен (6*) в данном случае можно написать в таком виде:

$$y \approx 9,22 - 6,90 \cdot \cos x + 3,45 \cdot \cos 2x + 3,30 \cdot \cos 3x - \\ - 0,62 \cdot \cos 4x + 0,60 \cos 5x + 0,25 \cdot \cos 6x + 21,09 \cdot \sin x - \\ - 2,80 \cdot \sin 2x + 1,97 \cdot \sin 3x + 1,07 \cdot \sin 4x - 1,02 \cdot \sin 5x. \quad (9)$$

Коэффициенты при всех членах этого ряда следует вычислять с одним и тем же числом десятичных знаков. При вычислении этих коэффициентов следует брать один лишний знак по сравнению с точностью исходных данных. Это делается потому, что при вычислении значений y коэффициенты a и b приходится умножать на косинусы и синусы соответствующих значений x , а затем брать сумму этих произведений. Эти вычисления как промежуточные необходимо производить с одним лишним знаком, который следует отбросить в окончательном результате (стр. 77).

Если в уравнении (9) положить $x = 0$, т. е. подставить значение x_0 , то мы находим y_0 , т. е. ординату точки пересечения кривой с осью ординат. Эта величина даётся первым основным уравнением или первой контрольной формулой и в данном случае оказывается равной 9,30.

Для того чтобы определить амплитуды и начальные фазы тех пяти гармонических составляющих, которые мы вычислили, и найти их уравнения, следует воспользоваться формулами (5'), подставляя в них соответственные значения a и b и принимая во внимание знаки каждого из a и b в отдельности.

Последнее необходимо для правильного определения углов φ . Так, углы φ_1 , φ_2 , φ_4 и φ_5 имеют отрицательные значения тангенсов, но знаки a и b показывают, что углы φ_1 и φ_4 лежат в четвёртой четверти (знаки a_1 и a_4 отрицательны, знаки b_1 и b_4 положительны), углы же φ_2 и φ_5 лежат, очевидно, во второй четверти (знаки a_2 и a_5 положительны, знаки b_2 и b_5 отрицательны); что же касается угла φ_3 , то он, очевидно, лежит в первой четверти, так как a_3 и b_3 оба положительны.

В результате таких вычислений находим для этих пяти гармонических составляющих такие значения A и φ :

$$A_1 = \sqrt{a_1^2 + b_1^2} = \sqrt{(-6,90)^2 + (21,09)^2} = 22,19,$$

$$\operatorname{tg} \varphi_1 = \frac{a_1}{b_1} = \frac{-6,90}{21,09} = -0,3272; \quad \varphi_1 = -18^\circ 5',$$

$$A_2 = \sqrt{a_2^2 + b_2^2} = \sqrt{(3,45)^2 + (-2,80)^2} = 4,44,$$

$$\operatorname{tg} \varphi_2 = \frac{a_2}{b_2} = \frac{3,45}{-2,80} = -1,2321; \quad \varphi_2 = 129^\circ 4',$$

$$A_3 = \sqrt{a_3^2 + b_3^2} = \sqrt{(3,30)^2 + (1,97)^2} = 3,86,$$

$$\operatorname{tg} \varphi_3 = \frac{a_3}{b_3} = \frac{3,30}{1,97} = 1,6751; \quad \varphi_3 = 59^\circ 10',$$

$$A_4 = \sqrt{a_4^2 + b_4^2} = \sqrt{(-0,62)^2 + (1,07)^2} = 1,23,$$

$$\operatorname{tg} \varphi_4 = \frac{-0,62}{1,07} = -0,5794; \quad \varphi_4 = -30^\circ 5',$$

$$A_5 = \sqrt{a_5^2 + b_5^2} = \sqrt{(0,60)^2 + (-1,02)^2} = 1,18,$$

$$\operatorname{tg} \varphi_5 = \frac{0,60}{-1,02} = -0,5882; \quad \varphi_5 = 149^\circ 32'.$$

Мы видим, что результат разложения функции y в гармонический ряд можно написать ещё в таком виде:

$$y \approx 9,22 + 22,19 \cdot \sin(x - 18^\circ 5') + 4,44 \cdot \sin(2x + 129^\circ 4') + \\ + 3,86 \cdot \sin(3x + 59^\circ 10') + 1,23 \cdot \sin(4x - 30^\circ 5') + \\ + 1,18 \cdot \sin(5x + 149^\circ 32') + 0,25 \cos 6x. \quad (9')$$

Периоды этих пяти гармонических составляющих мы

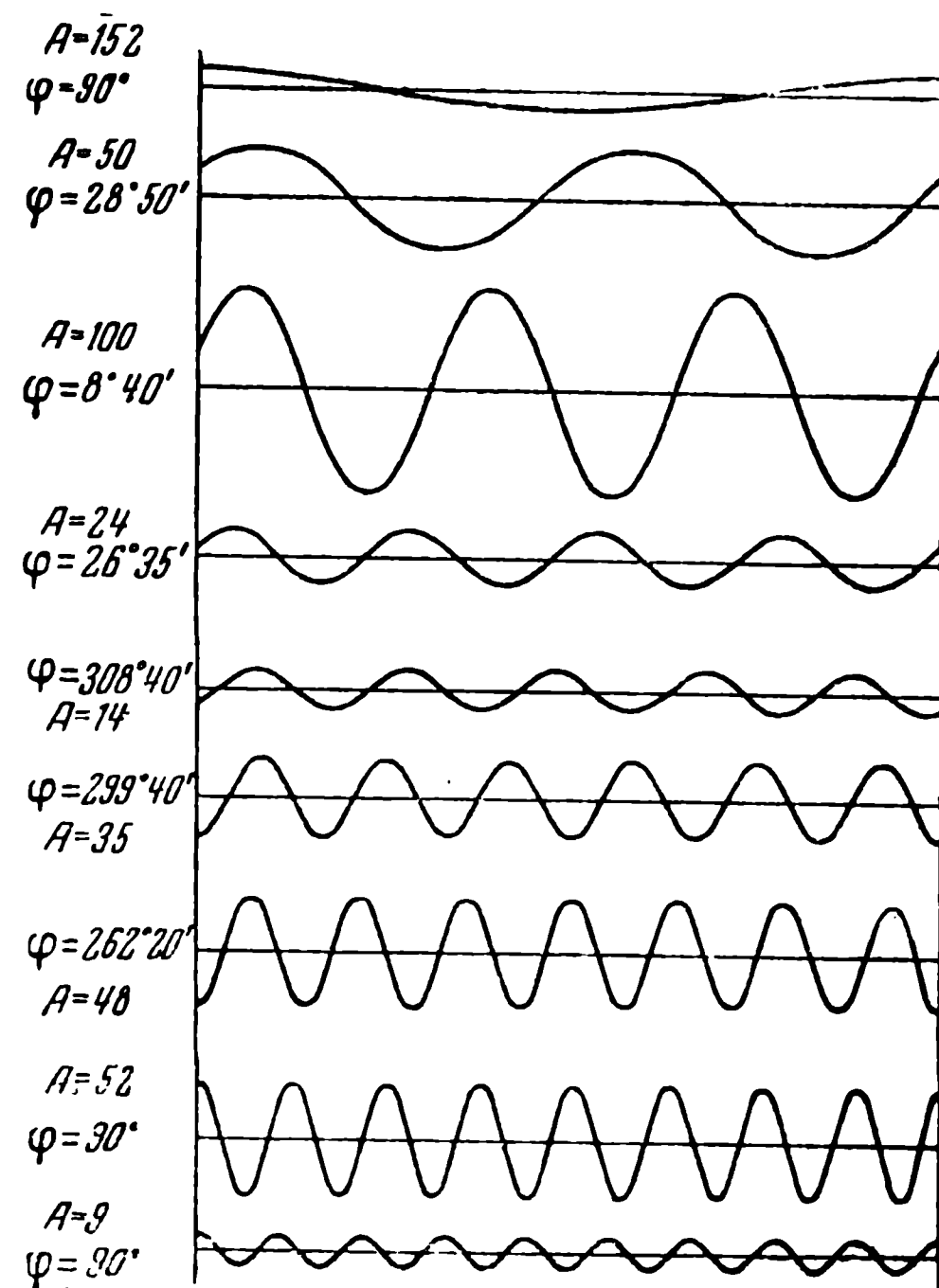


Рис. 50.

находим по условию (1), принимая во внимание период функции y , равный 2π . Имеем:

$$2\pi = T_0 = 2T_1 = 3T_2 = 4T_3 = 5T_4 \dots$$

Мы видим, что мы можем построить график функции y и её пяти гармонических составляющих. Период шестой гармонической составляющей также известен, но её амплитуду и начальную фазу мы определить не можем, так как эта составляющая представлена только одним последним членом в разложении (9'). Аналогичное разложение некоторой периодической функции на её гармонические составляющие дано на рис. 50, на котором изображены её девять первых гармоник.

Пример 2. Дан график периодической функции с периодом, равным 2π , изображённый на рис. 51. Найдём первые пять членов её разложения в ряд Фурье. Для

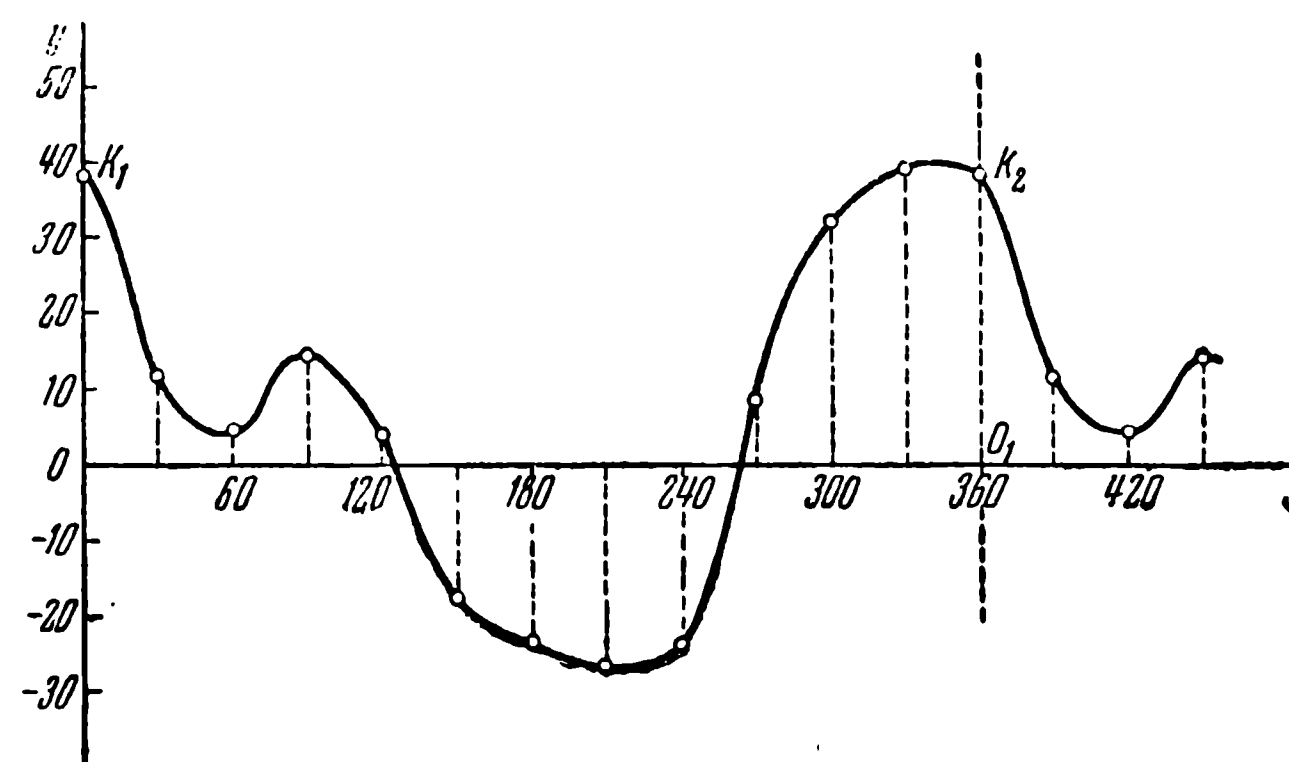


Рис. 51.

этого берём на кривой две точки K_1 и K_2 , которые находятся в одинаковых фазах, т. е. отстоят одна от другой на расстоянии одного периода. Выбор первой точки K_1 совершенно произволен, что же касается второй точки, то её следует определить возможно точнее. Для этого измеряют при помощи циркуля и миллиметровой линейки ординату точки K_1 и затем, отступив по кривой на расстояние, равное приблизительно периоду, находят на ней такую точку K_2 , ордината которой равна ординате точки K_1 . Точку O принимаем за начало координат и полный период кривой, т. е. расстояние OO_1

делим затем на 12 равных частей. В точках деления строим ординаты, длину которых измеряем также при помощи циркуля и линейки.

Допустим, что измерения дали результаты, указанные в таблице XXXI, где даны также соответствующие значения x .

Таблица XXXI

x	y	x	y
0°	38,2	180°	-23,3
30	12,0	210	-27,1
60	4,2	240	-24,2
90	14,4	270	8,1
120	4,1	300	32,3
150	-18,6	330	38,4

После этого мы производим все вычисления по данной выше схеме:

(y)	38,2	12,0	4,2	14,4	4,1	-18,6
	-23,3	38,4	32,3	8,1	-24,2	-27,1
Суммы (u)	14,9	50,4	36,5	22,5	-20,1	-45,7
Разности (v)	61,5	-26,4	-28,1	6,3	28,3	8,5
(u)	14,9	50,4	36,5			
	22,5	-45,7	-20,1			
Суммы (r)	37,4	4,7	16,4			
Разности (s)	-7,6	96,1	56,6			
(v)	-26,4	-28,1				
	8,5	28,3				
Суммы (p)	-17,9	0,2				
Разности (q)	-34,9	-56,4				
(r)	4,7	(q)	-34,9			
	16,4		-56,4			
(l)	21,1	(d)	-91,3			
(m)	-11,7	(h)	21,5			

Находим коэффициенты a и b :

$$a_0 = \frac{1}{12}(37,4 + 21,1) = 4,87,$$

$$a_1 = \frac{1}{6} \cdot \left(61,5 + \frac{3}{2} \cdot 96,1 + \frac{1}{2} \cdot 56,6 \right) = 28,84,$$

$$a_2 = \frac{1}{6} \cdot \left(-7,6 - \frac{1}{2} \cdot 11,7 \right) = -2,24,$$

$$a_3 = \frac{1}{6} \cdot (61,5 - 56,6) = 0,82,$$

$$a_4 = \frac{1}{6} \cdot \left(37,4 - \frac{1}{2} \cdot 21,1 \right) = 4,47,$$

$$a_5 = \frac{1}{6} \cdot \left(61,5 - \frac{\sqrt{3}}{2} \cdot 96,1 + \frac{1}{2} \cdot 56,6 \right) = 1,10,$$

$$a_6 = \frac{1}{12} \cdot (-7,6 + 11,7) = 0,34,$$

$$b_1 = \frac{1}{6} \cdot \left(6,3 - \frac{1}{2} \cdot 17,9 + \frac{3}{2} \cdot 0,2 \right) = -0,41,$$

$$b_2 = \frac{\sqrt{3}}{12} \cdot (-91,3) = -13,18,$$

$$b_3 = \frac{1}{6} \cdot (-17,9 - 6,3) = -4,03,$$

$$b_4 = \frac{\sqrt{3}}{12} \cdot 21,5 = 3,10,$$

$$b_5 = \frac{1}{6} \cdot \left(6,3 - \frac{1}{2} \cdot 17,9 - \frac{\sqrt{3}}{2} \cdot 0,2 \right) = -0,47.$$

Применяем контрольные формулы:

$$y_0 = 38,2 = 4,87 + 28,84 - 2,24 + 0,82 + 4,47 + \\ + 1,10 + 0,34 = 38,20,$$

$$v_1 = -26,4 = (-0,41 - 0,47) - 2 \cdot 4,03 + \\ + \sqrt{3}(-13,18 + 3,10) = -26,40.$$

Проверка показывает, что определение коэффициентов a и b сделано правильно. Таким образом, первые двенадцать членов разложения данной периодической функции

можно представить в таком виде:

$$y \approx 4,87 + 28,84 \cos x - 2,24 \cos 2x + 0,82 \cos 3x + \\ + 4,47 \cos 4x + 1,10 \cos 5x + 0,34 \cos 6x - 0,41 \sin x - \\ - 13,18 \sin 2x - 4,03 \sin 3x + 3,10 \sin 4x - 0,47 \sin 5x.$$

Тем же приёмом, который был разобран в первом примере, можно определить амплитуды и начальные фазы первых пяти гармонических составляющих и, если необходимо, дать затем графическую иллюстрацию этому разложению.

5. Технические приёмы гармонического анализа. Из той вычислительной схемы, которая была разобрана, и из примеров её применения мы видим, что анализ периодических кривых требует больших вычислений. Вычисления оказываются достаточно сложными даже при очень невысокой степени приближения, которую мы получили, определив первые пять гармонических составляющих. Если же мы, не довольствуясь этой степенью приближения, решаем увеличить число гармонических составляющих, например, до одиннадцати, то нам придётся увеличить число ординат до двадцати четырёх. В результате этого вычислительная схема значительно осложняется, и нам придётся выполнить очень большую вычислительную работу. Ещё более сложными оказываются вычисления, если мы применяем общий метод определения коэффициентов a и b , т. е. переходим к вычислению соответствующих интегралов (стр. 301). Эти интегралы в большинстве случаев допускают только приближённое вычисление при помощи одного из методов, разработанных для подобных операций.

Всё это было причиной того, что ещё очень давно начали разрабатывать различные искусственные, технические приёмы гармонического анализа, т. е. чисто технические приёмы приближённого определения коэффициентов a и b в разложении (4). Такие технические приёмы гармонического анализа можно разделить на две категории: во-первых, табличные приёмы гармонического анализа очень простые, но ограниченные в своих возможностях, и,

во-вторых, механические приёмы гармонического анализа, основанные на применении различных приборов, в том числе и очень сложных электрических приборов, к которым начинают переходить в последнее время.

1. **Табличные приёмы гармонического анализа.** Возможность применять таблицы при определении коэффициентов a и b основана на следующем:

во-первых, значения числовых коэффициентов при a и b во всех основных уравнениях повторяются: так, для случая двенадцати ординат, разобранного нами, коэффициентами при a и b во всех основных уравнениях

служили величины $0, \pm 1, \pm \frac{1}{2}, \pm \frac{\sqrt{3}}{2}$ и,

во-вторых, все основные уравнения составляются по вполне определённой схеме, т. е. коэффициенты a и b входят в них в одной и той же вполне определённой последовательности.

Табличный гармонический анализ наиболее разработан для случая 24 ординат. Это даёт возможность определять первые одиннадцать (двенадцать) гармонических составляющих. Значения x в этом случае изменяются, начиная от нуля, через каждые пятнадцать градусов, и соответственно этому синусы и косинусы в разложении (4) получают следующие значения:

$$\pm 1; \pm 0,966; \pm 0,866; \pm 0,707; \pm 0,500; \pm 0,259.$$

При вычислении коэффициентов a и b на эти числа приходится умножать значения всех двадцати четырёх ординат $y_1, y_2, y_3, \dots, y_{24}$. Для того чтобы не производить этих умножений для каждого разложения в отдельности, составлены произведения указанных значений синусов и косинусов на все числа от 1 до 200, вычисленные с достаточной для практических целей точностью. Таким образом, произведения, необходимые для определения коэффициентов a и b , мы находим непосредственно из таблиц.

Но эти произведения при определении каждого из a и b в отдельности необходимо составлять и подбирать в определённом числе и в определённой последовательности, как это можно видеть из основных формул. Для

того чтобы выбор этих произведений можно было производить возможно проще, чисто механически, разработаны особые таблицы и шаблоны, которые имеются в готовом виде. Эти таблицы и шаблоны составляются следующим образом:

1) На листах прозрачной бумаги, разделённой на удлиненные клетки (рис. 52), в первой горизонтальной

y	+1,000y	+0,986y	+0,866y	+0,707y	+0,500y	+0,259y	0	-0,259y	-0,500y	-0,707y	-0,866y	-0,986y	-1,000y
0													
1													
2													
3													
4													
5													
6													
7													
8													
9													
10													
11													
12													
13													
14													
15													
16													
17													
18													
19													
20													
21													
22													
23													
24													
	+							-					

Рис. 52.

строке дан указанный выше ряд значений синусов и косинусов, положительных на левой половине и отрицательных на правой. По вертикали в этой таблице (сё обыкновенно называют основной таблицей) имеется двадцать четыре строки, на которых следует выписать значения всех ординат, от первой до последней; для этого

служит первый вертикальный столбец таблицы со свободными клетками и пометкой сверху $+1,000y$. В остальных частях таблицы большинство клеток зачернено и только некоторые клетки остаются свободными. Эти свободные клетки отвечают тем произведениям, которые входят в основные уравнения. Так, коэффициент $+0,966$ приходится умножать на $y_1, y_5, y_7, y_{11}, y_{13}, y_{17}, y_{19}$ и y_{23} и только соответствующие клетки оставлены свободными; точно так же коэффициент $+0,707$ приходится умножать только на нечётные ординаты $y_1, y_3, y_5, y_7, y_9, y_{11}, y_{13}, y_{15}, y_{17}, y_{19}, y_{21}, y_{23}$, и только эти клетки оставлены свободными и т. д. Все свободные клетки в основной таблице следует заполнить соответствующими произведениями, которые берут из числовых таблиц. Таким образом, основную таблицу мы заполняем всеми данными, необходимыми для определения коэффициентов a и b , и это не требует от нас никакой вычислительной работы.

2) Для того чтобы производить из этих данных выбор произведений, необходимых для определения каждого из коэффициентов a и b в отдельности, применяются особые шаблоны такого же размера, как и основная таблица, и с такими же клетками. Каждый шаблон назначается для выбора произведений, необходимых при нахождении одного определённого коэффициента a или b . Таким образом, всего должно быть двадцать два шаблона по числу определяемых коэффициентов. Из этого числа одиннадцать шаблонов служат для определения коэффициентов $a_1, a_2, a_3, \dots, a_{10}$, т. е. для выбора членов ряда, дающего косинусональные компоненты, и одиннадцать шаблонов служат для определения коэффициентов $b_1, b_2, b_3, \dots, b_{10}$, т. е. для выбора членов ряда, дающего синусональные компоненты. На каждом из этих шаблонов имеется соответствующая пометка и нанесены красные клетки против тех мест таблицы, где находятся произведения, которые необходимо брать при определении данного коэффициента. Так, на рис. 53 дано схематическое изображение шаблона, который служит для определения коэффициента b_1 ; на рис. 53 соответствующие клетки зачернены. Накладывая на этот шаблон основную таблицу и совмещая

их клетки, мы выписываем произведения, которые оказываются на красном фоне, в отдельности те, которые приходятся на левой, и те, которые приходятся на правой половине таблицы, т. е. в отдельности положительные и отрицательные произведения. Их алгебраическая сумма, разделённая на двенадцать, даёт нам величину соответствующего коэффициента a или b . Такая операция, проведённая последовательно для всех шаблонов, даёт нам

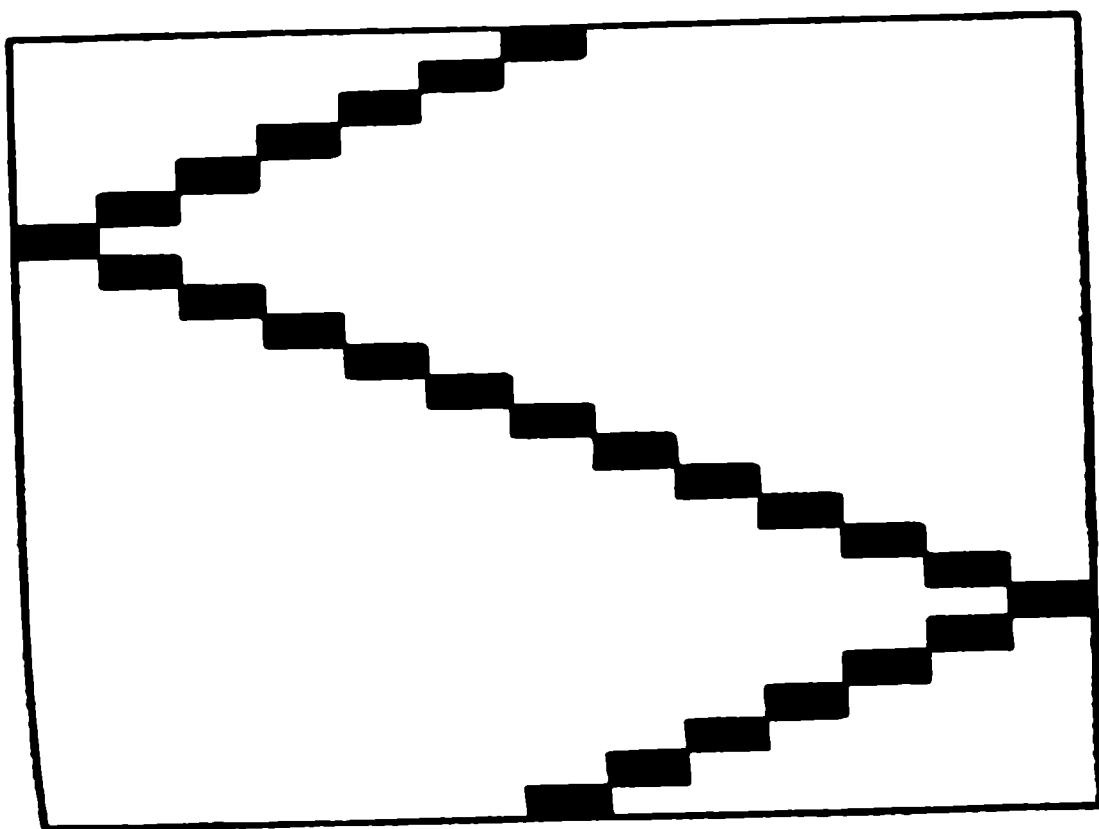


Рис. 53.

значения двадцати двух коэффициентов разложения (4). Что же касается свободного члена a_0 , то он попрежнему равен среднему арифметическому всех двадцати четырёх ординат и должен быть вычислен отдельно.

Таким образом, мы видим, что этот приём гармонического анализа требует очень небольшой вычислительной работы, которая ограничивается следующим:

1) Составив основную таблицу, необходимо выписать из неё двадцать два ряда слагаемых, определить алгебраическую сумму каждого ряда в отдельности и каждую из этих сумм разделить на двенадцать.

2) Необходимо определить среднее арифметическое всех двадцати четырёх ординат.

Следует также заметить, что для каждого разложения приходится составлять особую основ-

ную таблицу, что же касается шаблонов, то они остаются пригодными при всех разложениях. Наконец, необходимо ещё указать, что точность, которую мы получаем, пользуясь этим приёмом, в большинстве случаев оказывается вполне достаточной.

2. Механические приёмы гармонического анализа. Возможность определения коэффициентов a и b в разложении (4) механическим путём, т. е. при помощи различных приборов, в основном определяется тем, что эти коэффициенты оказываются равными значениям определённых интегралов (стр. 302), которые геометрически представляют собой не что иное, как площади, ограниченные осью абсцисс, двумя ординатами и отрезками соответствующих кривых. Таким образом, определяя эти площади, мы находим значения соответствующих коэффициентов a и b . На этом общем принципе основаны многочисленные приборы, известные под названием гармонических анализаторов. Такие приборы были разработаны рядом учёных. Не останавливаясь на описании этих приборов, укажем, что они различаются не только своей точностью и конструктивными особенностями, но и теми принципами, на которых они основаны. Так, некоторые гармонические анализаторы основаны на чисто механических процессах, подобных тем, которые применяются при измерении площадей обычными планиметрами, но известны и такие гармонические анализаторы, в которых используются электрические явления. Совершенно исключительные возможности в этом направлении были обнаружены в последние годы, когда к вычислительным операциям удалось применить электронно-лучевые приборы. Эти новые вычислительные методы в применении к вопросам гармонического анализа имеют необыкновенно широкие перспективы. К сожалению, этот новый метод на данной стадии своего развития приводит к необходимости применять чрезвычайно сложные радиотехнические узлы, что сильно ограничивает возможность его широкого использования.

ГЛАВА X

ЭМПИРИЧЕСКИЕ ФОРМУЛЫ

1. Общие замечания. Устанавливая из опыта ряд соответственных значений двух каких-либо величин (стр. 257), мы можем на основании этих данных подобрать алгебраическое выражение, которое в пределах опыта с достаточной точностью определяет зависимость между этими величинами. Такие алгебраические выражения принято называть *эмпирическими формулами*. Таким образом, эмпирические формулы подбираются на основании результатов данных измерений и должны давать в пределах опыта вполне удовлетворительное совпадение с экспериментальными результатами, причём значения постоянных коэффициентов в этих формулах определяются также на основании результатов измерений.

Поясним это на нескольких примерах.

Пример 1. Определение электрического сопротивления R проводника при определённых температурах привело к результатам, указанным в таблице XXXII, где дан ряд температур, выраженных в градусах Цельсия, и полученные при этих температурах значения R , выраженные в омах.

Таблица XXXII

$t^{\circ}\text{C}$	19,1	25,0	30,1	36,0	40,0	45,1	50,0
$R_{\text{ом}}$	76,30	77,80	79,75	80,80	82,35	83,90	85,10

Обрабатывая эти результаты одним из методов, которые изложены ниже, находим, что в данном случае зависимость между R и t можно с достаточной точностью выразить линейной формулой вида

$$R = a + bt. \quad (1)$$

Значение коэффициентов a и b в этой формуле можно определить, подставляя в неё соответственные значения R и t . Используя этот приём находим, что можно принять:

$$a = 70,90 \quad \text{и} \quad b = 0,284.$$

Таким образом, результаты измерений в данном случае можно выразить такой эмпирической формулой:

$$R = 70,90 + 0,284t. \quad (1')$$

Определение R по этой формуле для любых температур, которые лежат в пределах измерений, даёт вполне удовлетворительные результаты.

Пример 2. Чрезвычайно тщательные измерения коэффициента термического расширения β ртути показали, что его величина несколько возрастает с повышением температуры t . Классические исследования этого вопроса, произведённые Д. И. Менделеевым, привели к выводу, что зависимость β от t можно выразить линейной формулой, вполне аналогичной формуле (1), причём значения a и b следует принять равными:

$$a = 0,0001801 \quad \text{и} \quad b = 0,00000002.$$

Таким образом, эмпирическая формула, данная Д. И. Менделеевым для определения β ртути как функции t , имеет такой вид:

$$\beta_t = 0,0001801 + 0,00000002t. \quad (1'')$$

Значения a и b были вычислены с большой точностью для температурного интервала 0 и 100°C , и в этом интервале формула Д. И. Менделеева даёт прекрасное совпадение с опытными данными.

Пример 3. Упругость p насыщенного пара различных жидкостей, как известно, возрастает с температурой t по

очень сложному закону, для которого до настоящего времени не удаётся подыскать точного аналитического выражения. Вследствие этого было предложено очень большое количество различных приближённых формул, большая часть которых является чисто эмпирическими. Была предложена, например, эмпирическая формула такого вида:

$$\lg p = a + b \cdot \alpha^t + c \cdot \beta^t. \quad (2)$$

В этой формуле $\lg p$ оказывается определённой функцией температуры t , причём в эту формулу входит пять постоянных коэффициентов a , b , c , α и β . Значения этих постоянных коэффициентов оказываются различными не только для различных жидкостей, но даже для одной и той же жидкости в различных температурных интервалах. Так, из очень точных измерений упругости насыщенных паров воды следует, что значение коэффициентов a , b , c , α и β в этом случае оказывается равным:

1) в температурном интервале от -10 до $+100^\circ\text{C}$:

$$a = 4,732507, \quad \lg \alpha = 0,007014,$$

$$b = 0,013749, \quad \lg \beta = 1,996705,$$

$$c = -4,101985;$$

2) в температурном интервале от 100 до 200°C :

$$a = 6,299880, \quad \lg \alpha = 1,985245,$$

$$b = -2,190434, \quad \lg \beta = 1,998242,$$

$$c = -5,015341.$$

Мы видим, что при переходе от одного температурного интервала к другому имеет место весьма значительное изменение всех пяти коэффициентов. Вследствие этого, подставляя их величину в формулу (2), мы получаем две различные эмпирические формулы, которые можно применять только в соответственном интервале температур. Следует обратить внимание на то, что температура 100°C является конечной температурой первого интервала и одновременно начальной температурой второго интервала. Вследствие этого значение $\lg p$ при 100°C можно вычислить

как из первой формулы, так и из второй. Если произвести такие вычисления, то получаются два значения $\lg p$ при 100°C , чрезвычайно близкие одно к другому.

Таким образом, мы видим, что эмпирические формулы всегда имеют ограниченную область применения, которая не должна выходить за пределы опытных данных. Тем не менее эмпирические формулы имеют очень большое значение и весьма широко применяются. Часто имеют место даже такие случаи, когда эмпирическая формула, удачно подобранная, представляет собой окончательное завершение многолетних лабораторных исследований.

Широкое применение эмпирических формул объясняется главным образом двумя причинами:

во-первых, точное аналитическое выражение зависимости между исследуемыми величинами может оставаться нам неизвестным и мы по необходимости должны ограничиваться приближёнными формулами эмпирического характера, и

во-вторых, когда точная функциональная зависимость определяется формулой, настолько сложной, что её непосредственное применение при вычислениях было бы затруднительным.

Эмпирические формулы могут быть очень разнообразными, так как при выборе аналитической зависимости между исследуемыми величинами мы не руководствуемся какими-либо физическими теориями и ставим только одно главное условие: возможно близкое соответствие значений, вычисляемых по формуле, с опытными данными. В эмпирические формулы мы можем вводить различное число постоянных параметров или коэффициентов, величину которых нам приходится устанавливать с большой точностью. Более удачными следует считать эмпирические формулы с небольшим числом постоянных коэффициентов, два или три, так как большое число коэффициентов вносит значительные осложнения как при определении числовых значений этих коэффициентов, так и при дальнейшем применении формулы.

Необходимо указать, что, несмотря на большое разнообразие эмпирических формул, всё же имеется один вид

эмпирических зависимостей, который находит очень широкое распространение: чрезвычайно часто применяются эмпирические формулы в виде многочленов, расположенных по восходящим степеням независимого переменного и одновременно линейных по отношению ко всем постоянным коэффициентам. Очевидно, что эмпирические формулы этого типа в общем виде можно представить таким выражением:

$$y = f(x) \approx B_0 + B_1x + B_2x^2 + B_3x^3 + \dots + B_mx^m, \quad (3)$$

где величины $B_0, B_1, B_2, B_3, \dots, B_m$ обозначают постоянные коэффициенты, подлежащие определению.

Простыми примерами таких эмпирических зависимостей могут служить формулы (1') и (1''), которые можно рассматривать как сумму двух первых членов ряда (3), т. е. свободного члена и члена с первой степенью независимого переменного. В более сложных случаях появляются вторые и более высокие степени x . В общем случае по отношению к выражению (3) мы можем считать, что оно представляет собой сумму первых $(m+1)$ членов разложения функции $f(x)$ в бесконечный степенной ряд по степеням x . При определённых условиях такое разложение возможно для всех непрерывных функций; отсюда следует, что для любого ряда экспериментальных данных можно подобрать эмпирическую формулу вида (3), если только непрерывность данного явления не вызывает сомнений. Необходимо при этом иметь в виду, что точное аналитическое выражение функции $f(x)$ может оставаться нам неизвестным, и тем не менее оказывается возможным находить значения коэффициентов в её разложении.

Для определения коэффициентов B в уравнении (3) применяется несколько методов, из которых основными можно считать следующие четыре метода: 1) метод графический, 2) метод интерполяционных формул, 3) метод средних и 4) метод наименьших квадратов. Эти методы различаются как по точности, с которой можно определять коэффициенты B , так и по сложности вычислительных операций, которые при этом приходится производить. Вследствие этого рекомендуется выбирать тот или другой метод в зависимости от общих условий, имеющих место

в каждом отдельном случае. Очень часто удаётся получать достаточно точные значения при относительно несложных вычислениях, если применять два метода последовательно.

2. Графический метод. Если мы измеряем значения одной величины при определённых значениях другой величины, то неизбежным завершением такого ряда измерений является их графическое изображение обычно в прямоугольной системе координат. Это объясняется тем, что по виду кривых, получаемых на графиках, мы имеем возможность выяснить общий характер функциональной зависимости между исследуемыми величинами. Так, на графиках мы непосредственно видим, является ли данная функция возрастающей или убывающей и какому общему закону следует её изменение, т. е. изменяется ли она монотонно, имеет ли максимумы, минимумы или точки перегиба; так же ясно выявляется периодичность функций.

Общие приёмы и правила, которыми следует руководствоваться при построении точных графиков, уже были указаны (стр. 208). Но если мы применяем графический метод при подборе эмпирических формул, то необходимо сделать ещё одно весьма существенное разъяснение, которое заключается в следующем.

Как было указано, при построении графика следует проводить плавную кривую так, чтобы она проходила возможно ближе ко всем экспериментальным точкам. Однако это указание оставляет при построении кривых всё же очень большой произвол. Его можно частично устранить, воспользовавшись основным положением метода наименьших квадратов (стр. 358): сумма квадратов отклонений экспериментальных точек от кривой по вертикальному направлению, т. е. сумма квадратов величин e (рис. 54), должна быть наименьшей:

$$\sum e^2 = \text{минимум}. \quad (4)$$

Условие (4), вполне обоснованное теоретически, на практике не может вносить полной определённости, так как при построении кривых удовлетворить этому усло-

вию с достаточной строгостью является весьма сложной задачей. Вследствие этого прибегают к другому приёму, так называемому способу, или методу, выравнивания, который состоит в том, что кривую, построенную по экспериментальным данным, преобразуют в прямую линию (стр. 232).

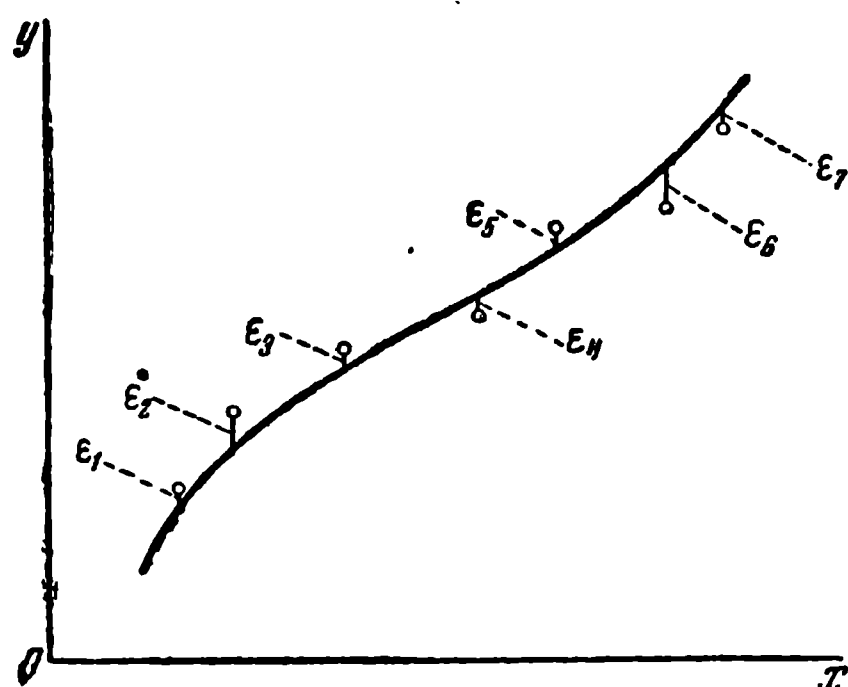


Рис. 54.

Если зависимость между y и x определяется некоторой кривой

$$y = f(x), \quad (5)$$

то для её преобразования в прямую линию надо ввести новые переменные X и Y , отвечающие условиям:

$$\begin{aligned} X &= \varphi_1(x, y), \\ Y &= \varphi_2(x, y), \end{aligned} \quad (6)$$

и подобрать функции φ_1 и φ_2 так, чтобы величины X и Y оказались связанными линейной зависимостью. В таком случае в системе координат (X, Y) мы получим некоторую прямую, уравнение которой в общей форме имеет такой вид:

$$Y = a + bX. \quad (7)$$

Величины X и Y мы определяем из уравнений (6) и строим соответствующую прямую возможно точнее, что в громадном большинстве случаев не представляет затруд-

нений, причём при проведении прямой значительно легче принять во внимание условие (4). Выбор функций φ_1 и φ_2 , необходимых для выпрямления кривой, чрезвычайно облегчается, если функция $f(x)$ нам известна или если нам известен хотя бы приближённый возможный вид функциональной зависимости между x и y . В последнем отношении очень большую помощь может оказать график кривой, непосредственно отвечающей экспериментальным данным, даже если этот график вычерчен только приближённо.

Построив прямую, мы можем определить и постоянные коэффициенты a и b в уравнении (7), так как значение a определяется ординатой точки пересечения этой прямой с осью Y , а значение b отвечает тангенсу угла, который она образует с осью X . Кроме того, мы можем подставить в уравнение (7) координаты двух каких-либо точек, через которые проходит эта прямая и которые нетрудно определить по графику. Из полученных таким приёмом двух уравнений мы можем также определить значения a и b , которые могут оказаться достаточно точными, в особенности если мы выбираем точки, расположенные далеко одна от другой, например вблизи концов прямой.

Определив значения коэффициентов a и b , мы получаем эмпирическую формулу в виде уравнения (7), которая устанавливает зависимость между величинами X и Y , т. е. между определёнными функциями величин x и y . На основании этой формулы можно затем установить приближённый вид функции $f(x)$, т. е. найти эмпирическую формулу, которая устанавливает зависимость непосредственно между величинами x и y .

На практике преобразование кривой в прямую линию часто достигается применением функциональных шкал и функциональных сеток (стр. 231), причём чрезвычайно широкое применение находят логарифмическая и полуллогарифмическая сетки. При применении функциональных координатных сеток с достаточно мелкими и точными делениями графический метод может дать хорошие результаты. Но вместе с тем необходимо иметь в виду то, что было сказано о графических вычислениях вообще (стр. 212), т. е. что эти методы не отличаются большой точностью.

Вследствие этого графический метод можно рекомендовать только в тех случаях, когда мы находим предварительные приближённые значения постоянных коэффициентов для эмпирических формул, или же в тех случаях, когда измерения выполнены с небольшой точностью, т. е. когда исходные данные имеют невысокую точность.

Для того чтобы яснее разобраться в вопросах применения графического метода при составлении эмпирических формул, рассмотрим несколько примеров.

Пример 1. Попробуем обосновать графическим методом эмпирическую формулу (1'), которая отвечает данным, приведённым в таблице XXXII. Если эти данные нанести графически в прямоугольной системе координат,

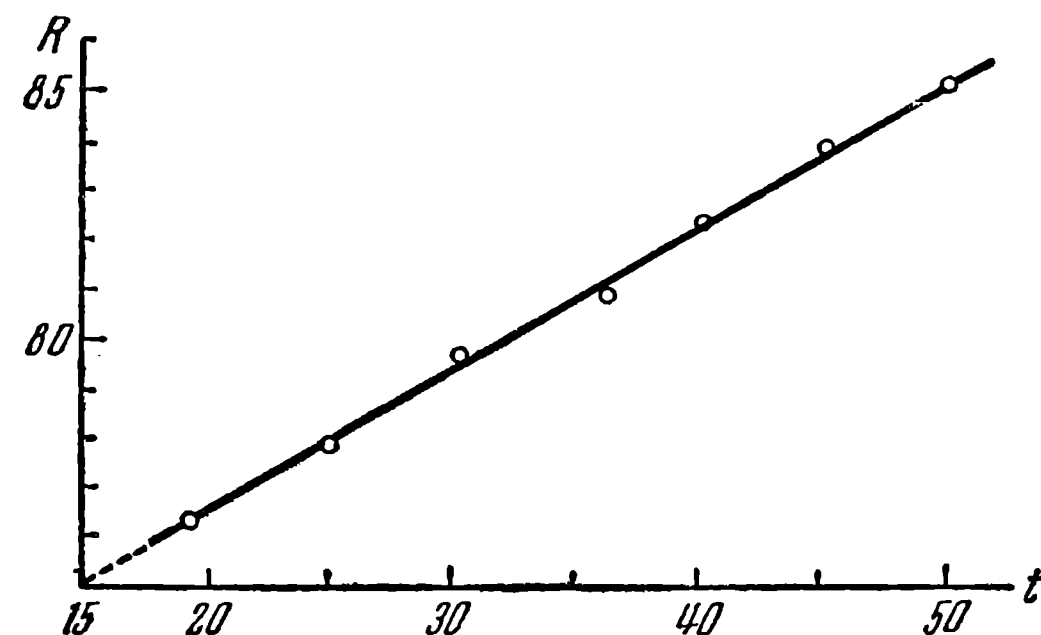


Рис. 55.

откладывая температуры t по оси абсцисс, а сопротивления R по оси ординат (рис. 55), то, как мы видим, экспериментальные точки очень хорошо ложатся на прямую линию. Таким образом, в данном случае линейной зависимостью связаны непосредственно величины R и t , что и дало возможность выразить их зависимость формулой (1).

Соответствующую прямую мы проводим так, чтобы было удовлетворено условие (4), по крайней мере насколько это возможно при построении точного графика, и выбираем на этой прямой две точки, достаточно удалённые. Одной из них мы можем считать точку с координатами

(50,0 и 85,10), т. е. последнюю из экспериментальных точек, которая очень хорошо ложится на прямую. Вторую точку следует выбрать вблизи другого конца прямой, например можем взять точку с абсциссой, равной 15; ординату этой точки примем равной 75,15. Подставляя координаты этих точек в уравнение (1), получаем два уравнения:

$$a + 50b = 85,10,$$

$$a + 15b = 75,15.$$

Определяя из этих уравнений значения a и b , находим:

$$b = \frac{9,95}{35} = 0,284 \quad \text{и} \quad a = 85,10 - 50 \cdot 0,284 = 70,90,$$

т. е. те значения коэффициентов, которые стоят в формуле (1'):

$$R = 70,90 + 0,284t.$$

Необходимо заметить, что если бы при вычислении значений a и b мы выбрали другие точки, то могли бы получить для a и b несколько иные значения. Так, например, мы могли бы взять точки с абсциссами 20 и 45; их ординаты можно принять равными соответственно 76,5 и 83,8. В этом случае получаются такие уравнения:

$$a + 45b = 83,8,$$

$$a + 20b = 76,5.$$

Отсюда находим такие значения a и b :

$$b = \frac{7,3}{25} = 0,292 \quad \text{и} \quad a = 83,8 - 45 \cdot 0,292 = 70,66.$$

Вследствие этого эмпирическая формула получает такой вид:

$$R = 70,66 + 0,292t. \quad (1'')$$

Отсюда следует, что точность, с которой мы в данном случае можем определить значения коэффициентов a и b , ограничена соответственно первым и вторым десятичными знаками. Но, несмотря на это, значения коэффициентов в эмпирических формулах принято давать с одним или двумя лишними десятичными знаками, которые следует отбрасывать только в окончательном результате (стр. 77).

Для проверки точности полученной эмпирической формулы, например формул (1') и (1''), следует подставить в них последовательно все данные значения t и, вычислив соответствующие значения R , сравнить их величину с экспериментальными данными. Такая проверка в данном случае приводит к результатам, которые даны в таблице XXXIII.

Таблица XXXIII

R экспериментальное	R , вычисленное по	
	формуле (1')	формуле (1'')
76,30	76,32	76,24
77,80	78,00	77,96
79,75	79,45	79,45
80,80	81,12	81,17
82,35	82,26	82,34
83,90	83,71	83,83
85,10	85,10	85,26

Хотя в этой таблице значения R , вычисленные по формулам (1') и (1''), приведены с двумя десятичными знаками, т. е. с той же точностью, как и в исходных данных, однако вычисленные значения R совпадают с его экспериментальными значениями только в целых единицах, а иногда в первом десятичном знаке. Эти результаты подтверждают указание на ограниченную точность графического метода.

Пример 2. Определение значений некоторой величины y при нескольких равноотстоящих значениях величины x привело к результатам, указанным в таблице XXXIV.

Таблица XXXIV

x	10	20	30	40	50	60	70	80
y	1,06	1,33	1,52	1,68	1,81	1,91	2,01	2,11

На основании этих результатов необходимо составить эмпирическую формулу, которой определялась бы зависимость y от x .

Для этого прежде всего табличные значения x и y изображаем графически в прямоугольной системе координат (рис. 56). Общий характер кривой, которая получается при этом, даёт основание предполагать, что в данном случае зависимость между y и x может отвечать степенной

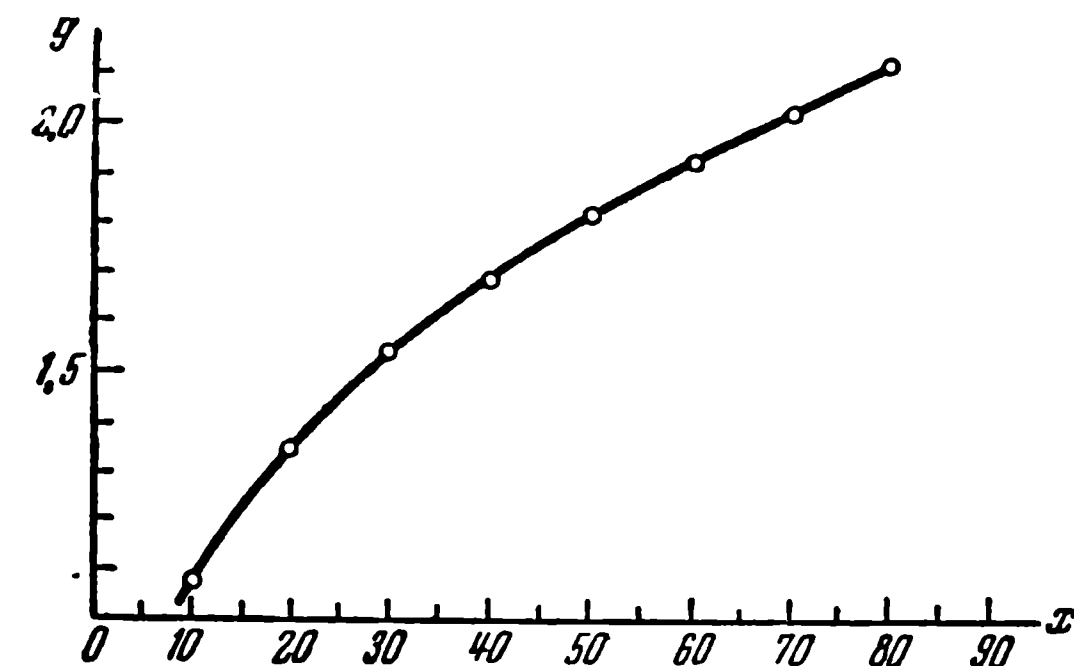


Рис. 56.

функции с дробным показателем степени (стр. 372). Вследствие этого допускаем, что в данном случае зависимость между y и x можно в общем виде представить выражением

$$y = ax^n, \quad (8)$$

где a и n обозначают некоторые постоянные параметры, величину которых следует определить на основании табличных данных.

Логарифмируя формулу (8), имеем:

$$\lg y = \lg a + n \lg x. \quad (8')$$

В этом выражении переходим к новым переменным, полагая:

$$\lg y = Y \quad \text{и} \quad \lg x = X. \quad (9)$$

На основании этих формул находим из выражения (8):

$$Y = \lg a + nX. \quad (10)$$

Мы получили уравнение, линейное относительно новых переменных Y и X , иначе говоря, если мы найдём из формул (9) значения Y и X и эти результаты изобразим графически, то должна получиться прямая линия. Вместо графического построения уравнения (10) в равномерной координатной сетке мы можем построить график уравнения (8') на логарифмической координатной сетке. Построив

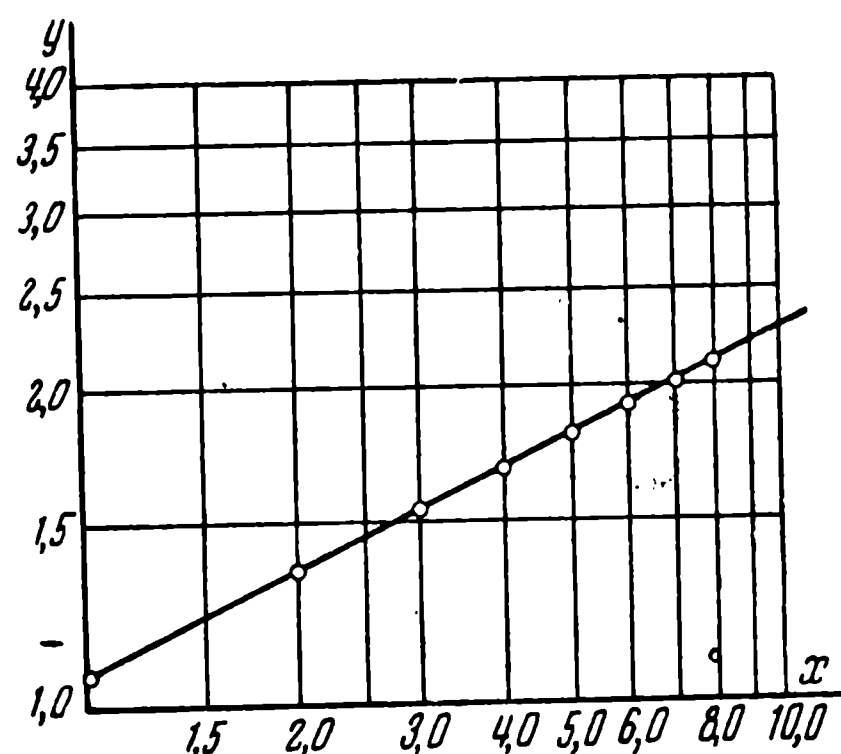


Рис. 57.

этот последний график (рис. 57), мы видим, что точки очень хорошо ложатся на прямую. Это доказывает правильность нашего основного положения, т. е. формулы (8). Вследствие этого переходим к определению параметров a и n в этой формуле. Для этого в уравнении (8') подставляем табличные значения координат первой и последней точек, как наиболее удалённых. Находим:

$$\lg 1,06 = \lg a + n \lg 10,$$

$$\lg 2,11 = \lg a + n \lg 80$$

или

$$\lg a + n = 0,0253,$$

$$\lg a + 1,9031n = 0,3243.$$

Отсюда получаем:

$$n = \frac{0,2990}{0,9031} = 0,3311,$$

$$\lg a = 0,0253 - 0,3311 = \bar{1},6942; \quad a = 0,4945.$$

Вставляя эти значения a и n в формулу (8), находим:

$$y = 0,4945x^{0,3311}. \quad (11)$$

Такой вид имеет эмпирическая формула, определяющая в данном случае зависимость y от x . Проверить пригодность и точность этой формулы можно тем же приёмом, который был указан в первом примере, т. е. в формулу (11) следует подставлять последовательно все табличные значения x и, вычисляя соответствующие значения y , сопоставлять их с его табличными значениями. Такая проверка в данном случае приводит ко вполне удовлетворительным результатам. Это объясняется главным образом тем, что при переходе к уравнению (10) мы получаем точки, которые очень точно ложатся на прямую линию.

Пример 3. Определение значений y для ряда равноотстоящих значений x привело к результатам, которые даны в таблице XXXV. Необходимо эти результаты представить в виде эмпирической формулы.

Таблица XXXV

x	1	2	3	4	5	6	7	8
y	15,3	20,5	27,4	36,6	49,1	65,6	87,8	117,6

Построив эти результаты графически (рис. 58), мы получаем кривую, типичную для показательных функций. Вследствие этого полагаем, что зависимость между x и y для данного случая можно в общем виде представить выражением

$$y = ae^{nx}, \quad (12)$$

где a и n обозначают постоянные параметры, которые необходимо определить.

Логарифмируя уравнение (12), находим:

$$\lg y = \lg a + n x \lg e$$

или

$$\lg y = \lg a + (n \lg e) \cdot x. \quad (12')$$

В этом выражении полагаем:

$$\lg y = Y. \quad (13)$$

Вследствие этого имеем:

$$Y = \lg a + (n \lg e) x. \quad (14)$$

Из этого уравнения, линейного по отношению к переменным Y и x , мы видим, что если воспользоваться

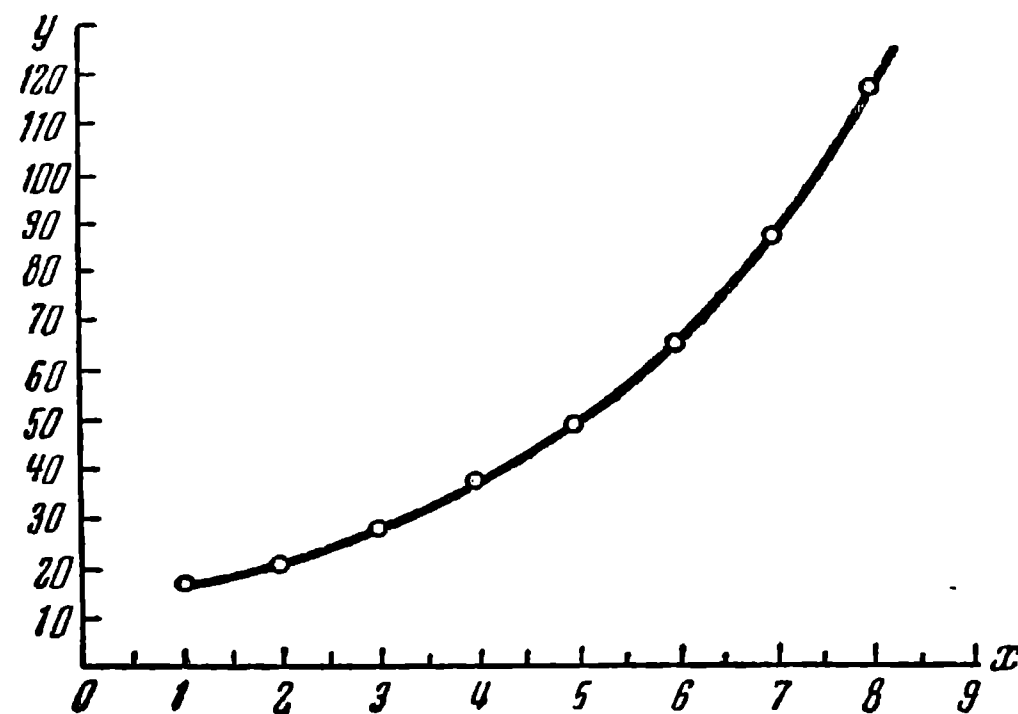


Рис. 58.

полулогарифмической координатной сеткой и откладывать значения x в равномерной шкале, а значения Y в логарифмической шкале, то мы должны получить прямую линию. Такое построение показывает, что все точки хорошо ложатся на прямую (рис. 59), что подтверждает наше основное предположение, т. е. формулу (12).

Подставляя в уравнение (12') табличные значения координат первой и последней точек, как наиболее удалённых, находим:

$$\begin{aligned} \lg 15,3 &= \lg a + n \lg e, \\ \lg 117,6 &= \lg a + 8n \lg e, \end{aligned}$$

или

$$\begin{aligned} \lg a + n \cdot \lg e &= 1,1847, \\ \lg a + 8n \cdot \lg e &= 2,0704. \end{aligned}$$

Отсюда находим:

$$n = \frac{0,8857}{7 \lg e} = 0,2913,$$

$$\lg a = 1,1847 - 0,1264 = 1,0583; \quad a = 11,44.$$

Таким образом, эмпирическая формула в данном случае имеет такой вид:

$$y = 11,44 \cdot e^{0,2913x}. \quad (15)$$

Пригодность и точность этой формулы можно проверить тем же приёмом, который был указан в первом и втором

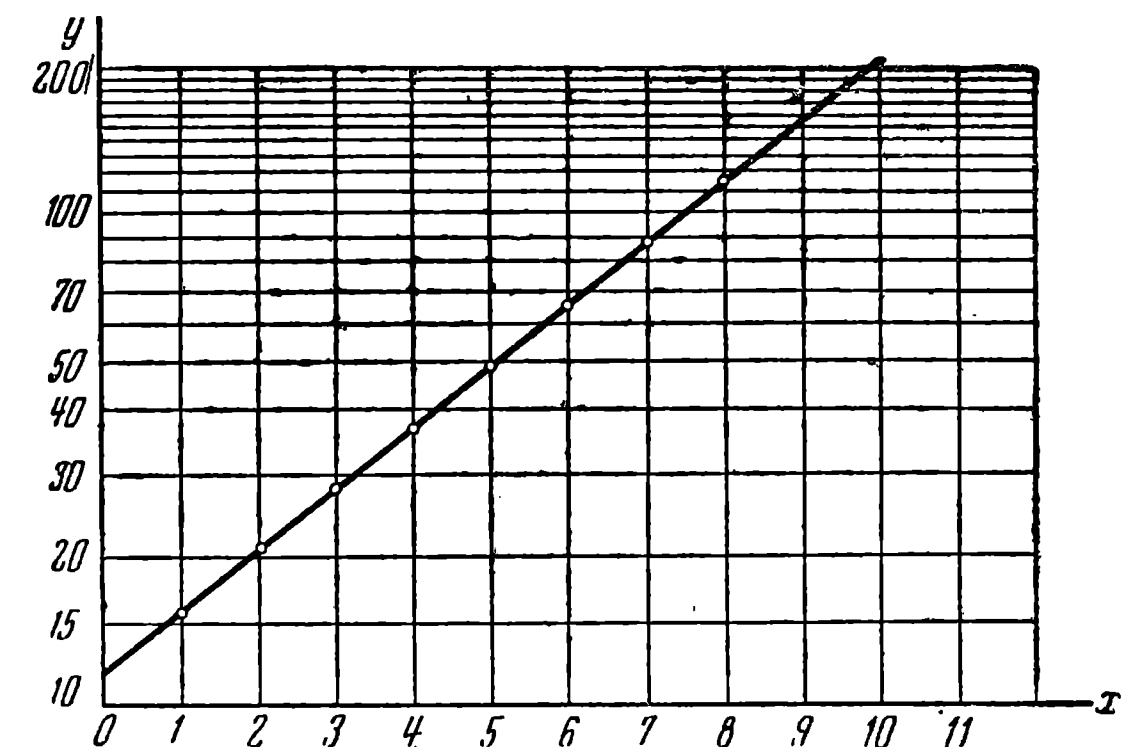


Рис. 59.

примерах. Такая проверка показывает, что значения y , вычисляемые по формуле (15), хорошо отвечают его табличным значениям, хотя всё же не в той мере, как во втором примере.

При применении логарифмической и полулогарифмической сеток следует иметь в виду, что началом координат в логарифмической сетке служит точка с координатами (1,1), а в полулогарифмической сетке точка с координатами (0,1), если равномерную шкалу мы строим на оси абсцисс.

3. Метод интерполяционных формул. Из того, что было сказано в первом параграфе, мы видим, что к числу эмпирических формул следует отнести и те интерполяционные формулы, которые мы подробно рассмотрели (гл. VIII). Действительно, каждую интерполяционную формулу мы берём с небольшим числом членов; такая формула всегда является приближённой и имеет область применения, ограниченную пределами измерений. Величину коэффициентов при последовательных членах интерполяционной формулы можно непосредственно вычислить, как это было показано.

Представляя интерполяционные формулы в виде ряда, расположенного по восходящим степеням x , и ограничиваясь его несколькими первыми членами, мы получаем эмпирическую формулу того вида, который в общей форме представлен выражением (3). Для того чтобы интерполяционные формулы привести к виду (3), их необходимо алгебраически преобразовать, причём следует воспользоваться тем видом интерполяционных формул, который представлен выражением (10) или первым из выражений (12), указанными в главе VIII. Такое преобразование мы уже выполнили (стр. 279), но только для случая, когда интерполяционная формула ограничена тремя членами, т. е. имеет вид:

$$y = y_0 + \frac{\Delta y_0}{h} (x - x_0) + \frac{\Delta^2 y_0}{2h^2} (x - x_0) (x - x_1).$$

В правой части этого выражения мы раскрыли скобки и, расположив полученный многочлен по восходящим степеням x , нашли, что

$$y = \left(y_0 - \frac{\Delta y_0}{h} x_0 + \frac{\Delta^2 y_0}{2h^2} x_0 x_1 \right) + \left(\frac{\Delta y_0}{h} - \frac{\Delta^2 y_0}{2h^2} x_0 - \frac{\Delta^2 y_0}{2h^2} x_1 \right) x + \frac{\Delta^2 y_0}{2h^2} x^2. \quad (16)$$

Сравнивая это выражение с общим выражением (3), мы видим, что коэффициенты при его последовательных членах, число которых в данном случае равно трём, т. е. коэффициенты B_0 , B_1 и B_2 , определяются такими

формулами:

$$\begin{aligned} B_0 &= y_0 - \frac{\Delta y_0}{h} x_0 + \frac{\Delta^2 y_0}{2h^2} x_0 x_1, \\ B_1 &= \frac{\Delta y_0}{h} - \frac{\Delta^2 y_0}{2h^2} x_0 - \frac{\Delta^2 y_0}{2h^2} x_1, \\ B_2 &= \frac{\Delta^2 y_0}{2h^2}. \end{aligned} \quad (17)$$

Если в интерполяционной формуле (10) (стр. 273) мы возьмём четыре первых члена, т. е. положим:

$$y = y_0 + \frac{\Delta y_0}{h} (x - x_0) + \frac{\Delta^2 y_0}{2h^2} (x - x_0) (x - x_1) + \frac{\Delta^3 y_0}{3! h^3} (x - x_0) (x - x_1) (x - x_2),$$

то совершенно тем же приёмом найдём, что

$$\begin{aligned} y = & \left(y_0 - \frac{\Delta y_0}{h} x_0 + \frac{\Delta^2 y_0}{2h^2} x_0 x_1 - \frac{\Delta^3 y_0}{3! h^3} x_0 x_1 x_2 \right) + \\ & + \left(\frac{\Delta y_0}{h} - \frac{\Delta^2 y_0}{2h^2} x_0 - \frac{\Delta^2 y_0}{2h^2} x_1 + \frac{\Delta^3 y_0}{3! h^3} x_0 x_1 + \right. \\ & + \left. \frac{\Delta^3 y_0}{3! h^3} x_0 x_2 + \frac{\Delta^3 y_0}{3! h^3} x_1 x_2 \right) x + \left(\frac{\Delta^2 y_0}{2h^2} - \frac{\Delta^3 y_0}{3! h^3} x_0 - \right. \\ & - \left. \frac{\Delta^3 y_0}{3! h^3} x_1 - \frac{\Delta^3 y_0}{3! h^3} x_2 \right) x^2 + \frac{\Delta^3 y_0}{3! h^3} x^3. \end{aligned}$$

Сравнивая это выражение с общим выражением (3), мы видим, что в этом случае следует взять четыре первых члена выражения (3), причём коэффициенты при этих членах, т. е. коэффициенты B_0 , B_1 , B_2 и B_3 , определяются такими формулами:

$$\begin{aligned} B_0 &= y_0 - \frac{\Delta y_0}{h} x_0 + \frac{\Delta^2 y_0}{2h^2} x_0 x_1 - \frac{\Delta^3 y_0}{3! h^3} x_0 x_1 x_2, \\ B_1 &= \frac{\Delta y_0}{h} - \frac{\Delta^2 y_0}{2h^2} x_0 - \frac{\Delta^2 y_0}{2h^2} x_1 + \frac{\Delta^3 y_0}{3! h^3} x_0 x_1 + \frac{\Delta^3 y_0}{3! h^3} x_0 x_2 + \frac{\Delta^3 y_0}{3! h^3} x_1 x_2, \\ B_2 &= \frac{\Delta^2 y_0}{2h^2} - \frac{\Delta^3 y_0}{3! h^3} x_0 - \frac{\Delta^3 y_0}{3! h^3} x_1 - \frac{\Delta^3 y_0}{3! h^3} x_2, \\ B_3 &= \frac{\Delta^3 y_0}{3! h^3}. \end{aligned} \quad (18)$$

Тем же приёмом можно получить аналогичные формулы и для тех случаев, когда мы берём ещё большее число членов в формуле (10), например пять, шесть и т. д. На этих формулах мы не останавливаемся подробнее ввиду их редкого применения.

Выражения (17) и (18) дают возможность находить значения коэффициентов B , так как все величины в правых частях этих выражений известны. Однако при вычислении коэффициентов B возникают некоторые осложнения, которые вызываются следующими причинами.

При определении значений y по интерполяционным формулам можно любую пару значений x и y считать начальными, как об этом было сказано раньше (стр. 278); действительно, мы убедились на примере, что при переходе от одной пары значений x и y к другой, результат вычислений не изменяется. В том примере, который мы разобрали при этом, были даны числа не приближённые, а точные; вследствие этого результаты вычислений во всех случаях оказались в точности одинаковыми. Но в формулы (17) и (18) входят величины y_0 и $x_0, x_1, x_2, \dots, x_n$, а также разности Δy_0 различных порядков; все эти величины являются не точными, а приближёнными. Вследствие этого, если при вычислении коэффициентов B мы будем применять в формулах (17) и (18) различные начальные данные, т. е. будем вводить в них различные пары значений x и y , то для каждого из коэффициентов при изменении начальных данных мы будем получать, вообще говоря, несколько различные значения. В результате мы встречаемся с необходимостью на основании этих различных значений коэффициентов B определить их наилучшее значение.

Рассмотрим эти вопросы на каком-либо примере. Возьмём хотя бы результаты исследования скорости той химической реакции, которую мы уже разбирали раньше (стр. 280). Мы нашли из таблицы разностей, что в данном случае вторые разности можно считать постоянными. На основании этого интерполяционную формулу для определения процентного содержания A вещества в системе в зависимости от времени t можно написать в та-

ком виде:

$$A = A_0 + \frac{\Delta A_0}{h} (t - t_0) + \frac{\Delta^2 A_0}{2h^2} (t - t_0) (t - t_1).$$

Таким образом, представляя эмпирическую формулу в виде многочлена, расположенного по восходящим степеням t , мы можем для данного случая в общем виде принять:

$$A = B_0 + B_1 t + B_2 t^2.$$

Для определения коэффициентов B следует применить формулы (17). Величины, которые входят в правые части этих формул, берём из исходных данных и из таблицы разностей, причём:

1) Величину h будем выражать попрежнему в минутах, т. е. примем, что

$$h = 5.$$

2) Для второй разности примем её среднее значение, полагая, что

$$\Delta^2 A = 1,14.$$

3) Что же касается величин A_0, t_0 и t_1 , а также величины ΔA_0 , то мы вычислим значения коэффициентов B , применяя последовательно все соответственные значения t и A , которыми мы располагаем в данном случае. Число этих значений t и A определяется, очевидно, числом значений первых разностей ΔA . Таким образом, мы можем ввести при вычислениях шесть значений t и A ; их последними значениями, т. е. $t = 37$ и $A = 36,3$, мы воспользоваться не можем, так как для этого значения A мы не имеем соответствующего значения ΔA .

Отсюда следует, что мы можем найти в данном случае шесть возможных значений коэффициентов B_0 и B_1 . Что же касается коэффициента B_2 , то его величина зависит только от $\Delta^2 A_0$ и h , т. е. при всех значениях t и A она остаётся постоянной.

Производим все эти вычисления.

1. Определяем значение B_2 . На основании третьей формулы (17) имеем:

$$B_2 = \frac{\Delta^2 A}{2h^2} = \frac{1,14}{2 \cdot 5^2} = 0,0228.$$

2. Определяем значения B_0 и B_1 . На основании первой и второй формул (5) имеем:

$$B_0 = A_0 - \frac{\Delta A_0}{h} t_0 + \frac{\Delta^2 A_0}{2h^2} t_0 \cdot t_1,$$

$$B_1 = \frac{\Delta A_0}{h} - \frac{\Delta^2 A_0}{2h^2} t_0 - \frac{\Delta^2 A_0}{2h^2} t_1.$$

В эти формулы подставляем последовательно все шесть возможных значений A_0 и t_0 и соответственные значения ΔA и t . Получаем такие результаты:

1) При

$$A_0 = 83,7, \quad t_0 = 7, \quad t_1 = 12 \quad \text{и} \quad \Delta A_0 = -10,8$$

имеем:

$$B_0 = 83,7 + \frac{10,8}{5} \cdot 7 + \frac{1,14}{2 \cdot 5^2} \cdot 7 \cdot 12 = 100,735,$$

$$B_1 = -\frac{10,8}{5} - \frac{1,14}{2 \cdot 5^2} \cdot 7 - \frac{1,14}{1 \cdot 5^2} \cdot 12 = -2,5932.$$

2) При

$$A_0 = 72,9, \quad t_0 = 12, \quad t_1 = 17 \quad \text{и} \quad \Delta A_0 = -9,7$$

имеем:

$$B_0 = 72,9 + \frac{9,7}{5} \cdot 12 + \frac{1,14}{2 \cdot 5^2} \cdot 12 \cdot 17 = 100,831,$$

$$B_1 = -\frac{9,7}{5} - \frac{1,14}{2 \cdot 5^2} \cdot 12 - \frac{1,14}{2 \cdot 5^2} \cdot 17 = -2,6012.$$

3) При

$$A_0 = 63,2, \quad t_0 = 17, \quad t_1 = 22 \quad \text{и} \quad \Delta A_0 = -8,5$$

имеем:

$$B_0 = 63,2 + \frac{8,5}{5} \cdot 17 + \frac{1,14}{2 \cdot 5^2} \cdot 17 \cdot 22 = 100,627,$$

$$B_1 = -\frac{8,5}{5} - \frac{1,14}{2 \cdot 5^2} \cdot 17 - \frac{1,14}{2 \cdot 5^2} \cdot 22 = -2,5892.$$

4) При

$$A_0 = 54,7, \quad t_0 = 22, \quad t_1 = 27 \quad \text{и} \quad \Delta A_0 = -7,2$$

имеем:

$$B_0 = 54,7 + \frac{7,2}{5} \cdot 22 + \frac{1,14}{2 \cdot 5^2} \cdot 22 \cdot 27 = 99,923,$$

$$B_1 = -\frac{7,2}{5} - \frac{1,14}{2 \cdot 5^2} \cdot 22 - \frac{1,14}{2 \cdot 5^2} \cdot 27 = -2,5572.$$

5) При

$$A_0 = 47,5, \quad t_0 = 27, \quad t_1 = 32 \quad \text{и} \quad \Delta A_0 = -6,1$$

имеем:

$$B_0 = 47,5 + \frac{6,1}{5} \cdot 27 + \frac{1,14}{2 \cdot 5^2} \cdot 27 \cdot 32 = 100,139,$$

$$B_1 = -\frac{6,1}{5} - \frac{1,14}{2 \cdot 5^2} \cdot 27 - \frac{1,14}{2 \cdot 5^2} \cdot 32 = -2,5652.$$

6) При

$$A_0 = 41,4, \quad t_0 = 32, \quad t_1 = 37 \quad \text{и} \quad \Delta A_0 = -5,1$$

имеем:

$$B_0 = 41,4 + \frac{5,1}{5} \cdot 32 + \frac{1,14}{2 \cdot 5^2} \cdot 32 \cdot 37 = 101,035,$$

$$B_1 = -\frac{5,1}{5} - \frac{1,14}{2 \cdot 5^2} \cdot 32 - \frac{1,14}{2 \cdot 5^2} \cdot 37 = -2,5932.$$

Для каждого из коэффициентов B_0 и B_1 мы получили по шести различных значений. Их среднее арифметическое можно принять как окончательное значение этих коэффициентов. Выполняя эти вычисления, находим:

$$B_0 = 100,55 \quad \text{и} \quad B_1 = -2,583.$$

Таким образом, эмпирическую формулу для данной химической реакции можно написать в таком виде:

$$A = 100,55 - 2,583 t + 0,0228 t^2. \quad (19)$$

Для проверки этой формулы следует подставить в её правую часть все данные значения t , т. е. 7, 12, 17, ..., и вычислить соответствующие значения A . Результаты таких вычислений приведены в таблице XXXVI, где дана также разность между значениями A , вычисленными по формуле (19), и его табличными значениями; эти разности обозначены через ϵ .

Таблица XXXVI

A вычисленное	A экспериментальное	ϵ
83,58	83,7	-0,12
72,83	72,9	-0,07
63,22	63,2	+0,02
54,75	54,7	+0,05
47,42	47,5	-0,08
41,23	41,4	-0,17

Из этой таблицы мы видим, что формула (19) даёт значения A , хотя и близкие к его табличным значениям, но всё же заметно от них отличающиеся: для большинства точек расхождение замечается уже в первом десятичном знаке, т. е. в том знаке, который в табличных данных является точным. Эти результаты дают основания считать, что точность определения коэффициентов B в разложении (3) при помощи интерполяционных формул является не особенно высокой. Кроме того, применение этого метода требует сложных вычислений, в особенности при малой сходимости ряда (3), когда приходится брать в нём большое число членов.

По отношению к формуле (19) необходимо сделать ещё одно замечание. Если бы мы эту формулу попробовали применить к начальному моменту времени, т. е. положили бы в ней

$$t = 0,$$

то нашли бы, что

$$A = 100,55.$$

Этот результат является, очевидно, совершенно неподобным, так как для начального момента A должно быть равно 100%, и значение A больше 100%, к которому приводит формула (19), невозможно. Однако это нельзя рассматривать как органический порок формулы (19): эту формулу, как и все эмпирические формулы, можно применять только в том

интервале значений x , для которого она рассчитана, в данном случае в пределах от семи до тридцати семи минут, и если мы выходим за пределы этого интервала, то всегда можем получить неправильный результат (стр. 293).

4. Метод средних. Если результаты измерений мы изображаем графически и затем проводим плавную кривую возможно ближе ко всем точкам (рис. 60), то по отношению

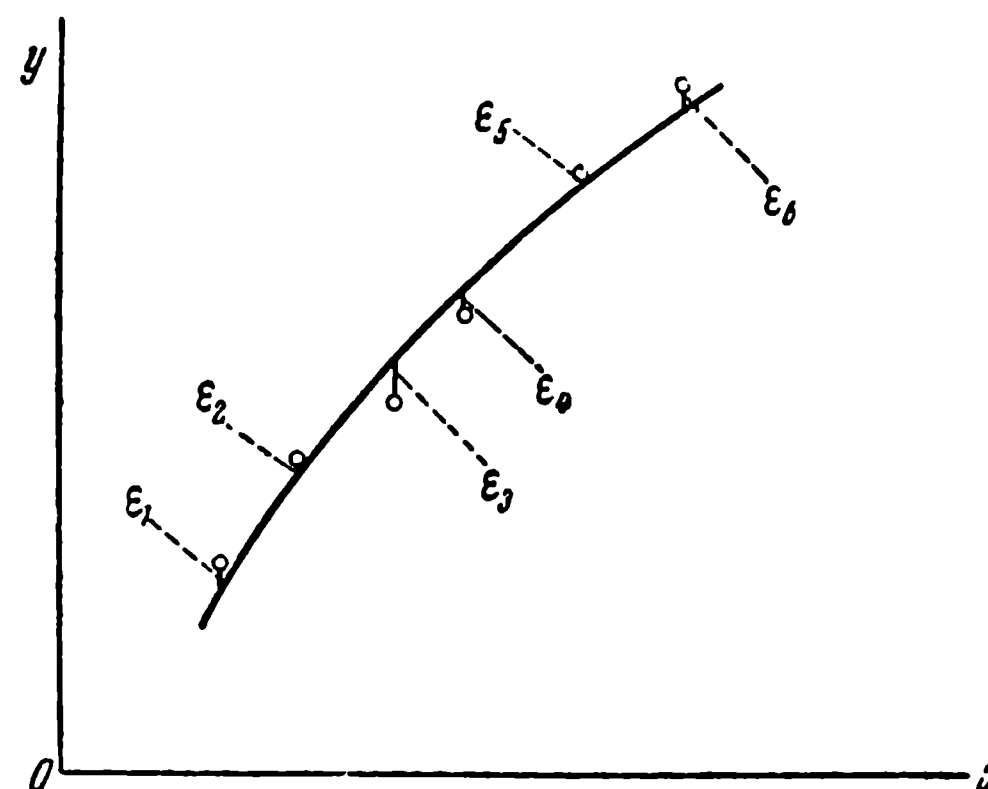


Рис. 60.

к отклонениям экспериментальных точек от кривой по вертикальному направлению, т. е. по отношению к величинам ϵ , мы можем сделать два основных предположения:

1) Во-первых, мы можем предположить, что сумма квадратов величин ϵ должна быть наименьшей, т. е. допустить, что имеет место условие (4), которое, как уже было сказано, является основным положением метода наименьших квадратов. Таким образом, здесь мы предполагаем, что лучшей кривой должна быть та кривая, которая удовлетворяет условию (4).

2) Во-вторых, мы можем ограничиться более простым предположением, которое заключается в следующем. Так как одни из ϵ будут иметь положительные знаки (экспериментальные точки лежат ниже кривой), а другие ϵ будут

иметь отрицательные знаки (экспериментальные точки лежат выше кривой), то мы можем условно принять, что лучшей кривой должна быть та, для которой алгебраическая сумма отклонений ϵ будет равна нулю, т. е. будем искать такую кривую, для которой

$$\sum \epsilon = 0. \quad (20)$$

Это условие даёт возможность обосновать очень простой метод определения коэффициентов B , который принято называть методом средних. Для того чтобы мы могли при помощи метода средних находить значения коэффициентов B в разложении (3), необходимо прежде всего определить число членов этого ряда, которое следует вводить при вычислениях. Для этого необходимо составить таблицу разностей (стр. 264) и определить порядок тех разностей, которые можно принять постоянными. В результате этого мы можем в разложении (3) ограничиться вполне определённым числом членов.

Затем мы все члены ряда (3) собираем в правой части. Таким образом, мы получаем выражение, правая часть которого в общей форме имеет такой вид:

$$B_0 + B_1x + B_2x^2 + B_3x^3 + \dots + B_mx^m - y. \quad (21)$$

В это выражение мы подставляем последовательно все значения x и y и приравниваем полученные выражения соответствующим отклонением ϵ . Если мы имеем в своём распоряжении n экспериментальных точек:

$$(x_1, y_1), (x_2, y_2), (x_3, y_3), \dots, (x_n, y_n),$$

то после их подстановки в выражение (21) мы получаем такую систему уравнений:

$$\begin{aligned} B_0 + B_1 x_1 + B_2 x_1^2 + B_3 x_1^3 + \dots + B_m x_1^m - y_1 &= \varepsilon_1, \\ B_0 + B_1 x_2 + B_2 x_2^2 + B_3 x_2^3 + \dots + B_m x_2^m - y_2 &= \varepsilon_2, \\ B_0 + B_1 x_3 + B_2 x_3^2 + B_3 x_3^3 + \dots + B_m x_3^m - y_3 &= \varepsilon_3, \\ &\vdots \\ &\vdots \\ B_0 + B_1 x_m + B_2 x_m^2 + B_3 x_m^3 + \dots + B_m x_m^m - y_m &= \varepsilon_m, \end{aligned} \quad (22)$$

где через $\varepsilon_1, \varepsilon_2, \varepsilon_3, \dots, \varepsilon_n$ обозначена величина соответствующих отклонений.

Таким образом, мы получили систему n уравнений с $(m+1)$ неизвестными $B_0, B_1, B_2, B_3, \dots, B_m$. Эти уравнения принято называть начальными уравнениями.

Так как число коэффициентов B , которые подлежат определению, обычно бывает невелико, например три, четыре, то число экспериментальных точек оказывается значительно больше числа коэффициентов B , иначе говоря, на практике n всегда оказывается больше $(m + 1)$. Вследствие этого мы можем все n начальных уравнений разбить на столько отдельных групп, сколько коэффициентов B необходимо определить. По отношению к этим группам уравнений мы делаем предположение, что для каждой группы в отдельности условие (20) остаётся в силе. Вследствие этого, если мы сложим все уравнения в каждой группе в отдельности, то правые части полученных выражений мы можем считать равными нулю. Таким приёмом мы уменьшаем число уравнений до числа неизвестных, причём ни одно из начальных уравнений не остаётся при этом неиспользованным. В результате из полученной системы $(m + 1)$ уравнений с $(m + 1)$ неизвестными мы определяем их значения одним из обычных алгебраических приёмов.

Нетрудно видеть, что все вычисления на практике должны быть значительно более простыми по сравнению с теми, которые были установлены нами при чисто теоретическом обосновании метода средних.

Для того чтобы лучше познакомиться с этим методом, рассмотрим в качестве примера ту химическую реакцию, которую мы уже разбирали. Составим для этой реакции эмпирическую формулу, пользуясь методом средних. Для этой реакции, как мы нашли, многочлен (3) можно представить в таком виде:

$$A = B_0 + B_1 t + B_2 t^2.$$

В нашем распоряжении имеется семь экспериментальных точек (стр. 280). Вследствие этого получаем семь началь-

ных уравнений, которые следует написать в таком виде:

$$\begin{aligned} 1) & B_0 + 7B_1 + 49B_2 - 83,7 = \varepsilon_1, \\ 2) & B_0 + 12B_1 + 144B_2 - 72,9 = \varepsilon_2, \\ 3) & B_0 + 17B_1 + 289B_2 - 63,2 = \varepsilon_3, \\ 4) & B_0 + 22B_1 + 484B_2 - 54,7 = \varepsilon_4, \\ 5) & B_0 + 27B_1 + 729B_2 - 47,5 = \varepsilon_5, \\ 6) & B_0 + 32B_1 + 1024B_2 - 41,4 = \varepsilon_6, \\ 7) & B_0 + 37B_1 + 1369B_2 - 36,3 = \varepsilon_7. \end{aligned} \quad (23)$$

Так как нам необходимо определить три коэффициента B_0 , B_1 и B_2 , то мы должны эти семь уравнений разбить на три группы, что, очевидно, можно сделать весьма различными способами, причём каждая группировка уравнений приводит к различным значениям коэффициентов B . Не представляется возможным определить заранее, какой способ группировки уравнений может обеспечить лучшие значения коэффициентов B . Тем не менее практика показывает, что наиболее удачные эмпирические формулы обычно получаются, если уравнения группировать в той последовательности, которая отвечает последовательности экспериментальных данных, и вводить в каждую группу, по возможности, одинаковое число начальных уравнений.

В данном случае, разбивая семь уравнений на три группы, мы должны взять в двух группах по два уравнения и в одной группе три уравнения. Если в эти группы вводить уравнения в той последовательности, в которой они даны, то, очевидно, мы можем образовать только три группировки уравнений такого вида:

$$\begin{aligned} \text{I. } 1) & B_0 + 7B_1 + 49B_2 - 83,7 = \varepsilon_1, & \text{Первая группа} \\ & B_0 + 12B_1 + 144B_2 - 72,9 = \varepsilon_2. \\ 2) & B_0 + 17B_1 + 289B_2 - 63,2 = \varepsilon_3, & \text{Вторая группа} \\ & B_0 + 22B_1 + 484B_2 - 54,7 = \varepsilon_4. \\ 3) & B_0 + 27B_1 + 729B_2 - 47,5 = \varepsilon_5, & \text{Третья группа} \\ & B_0 + 32B_1 + 1024B_2 - 41,4 = \varepsilon_6, \\ & B_0 + 37B_1 + 1369B_2 - 36,3 = \varepsilon_7. \end{aligned}$$

$$\begin{aligned} \text{II. } 1) & B_0 + 7B_1 + 49B_2 - 83,7 = \varepsilon_1, & \text{Первая группа} \\ & B_0 + 12B_1 + 144B_2 - 72,9 = \varepsilon_2. \\ 2) & B_0 + 17B_1 + 289B_2 - 63,2 = \varepsilon_3, & \text{Вторая группа} \\ & B_0 + 22B_1 + 484B_2 - 54,7 = \varepsilon_4, \\ & B_0 + 27B_1 + 729B_2 - 47,5 = \varepsilon_5. \\ 3) & B_0 + 32B_1 + 1024B_2 - 41,4 = \varepsilon_6, & \text{Третья группа} \\ & B_0 + 37B_1 + 1369B_2 - 36,3 = \varepsilon_7. \\ \text{III. } 1) & B_0 + 7B_1 + 49B_2 - 83,7 = \varepsilon_1, & \text{Первая группа} \\ & B_0 + 12B_1 + 144B_2 - 72,9 = \varepsilon_2, \\ & B_0 + 17B_1 + 289B_2 - 63,2 = \varepsilon_3. \\ 2) & B_0 + 22B_1 + 484B_2 - 54,7 = \varepsilon_4, & \text{Вторая группа} \\ & B_0 + 27B_1 + 729B_2 - 47,5 = \varepsilon_5. \\ 3) & B_0 + 32B_1 + 1024B_2 - 41,4 = \varepsilon_6, & \text{Третья группа} \\ & B_0 + 37B_1 + 1369B_2 - 36,3 = \varepsilon_7. \end{aligned}$$

Следует заметить также, что если бы мы при группировках уравнений не придерживались указанной последовательности и допускали возможность соединения уравнений в группах по принципу сочетаний, то число возможных группировок оказалось бы чрезвычайно большим. Это число N , которое в данном случае определяется формулой

$$N = \frac{7!}{2!2!3!},$$

было бы равно 210. Такое число группировок явно выходит за пределы возможного их использования при выборе лучшей эмпирической формулы.

Исследуем все три указанные выше группировки уравнений.

I. Первая группировка:

$$\begin{aligned} 1) & B_0 + 7B_1 + 49B_2 - 83,7 = \varepsilon_1, \\ & B_0 + 12B_1 + 144B_2 - 72,9 = \varepsilon_2. \\ 2) & B_0 + 17B_1 + 289B_2 - 63,2 = \varepsilon_3, \\ & B_0 + 22B_1 + 484B_2 - 54,7 = \varepsilon_4. \\ 3) & B_0 + 27B_1 + 729B_2 - 47,5 = \varepsilon_5, \\ & B_0 + 32B_1 + 1024B_2 - 41,4 = \varepsilon_6, \\ & B_0 + 37B_1 + 1369B_2 - 36,3 = \varepsilon_7. \end{aligned}$$

Складывая уравнения в каждой из этих групп и приравнявая эти суммы нулю, получаем такие три уравнения:

$$\begin{aligned} 2B_0 + 19B_1 + 193B_2 &= 156,6, \\ 2B_0 + 39B_1 + 773B_2 &= 117,9, \\ 3B_0 + 96B_1 + 3122B_2 &= 125,2. \end{aligned}$$

Определяя из этих уравнений B_0 , B_1 и B_2 , находим такие результаты:

$$\begin{aligned} B_0 &= 100,9606, \\ B_1 &= -2,6281, \\ B_2 &= 0,0239. \end{aligned}$$

Такие значения коэффициентов B_0 , B_1 и B_2 даёт первая группировка уравнений.

II. Вторая группировка:

$$\begin{aligned} 1) \quad & B_0 + 7B_1 + 49B_2 - 83,7 = \varepsilon_1, \\ & B_0 + 12B_1 + 144B_2 - 72,9 = \varepsilon_2, \\ 2) \quad & B_0 + 17B_1 + 289B_2 - 63,2 = \varepsilon_3, \\ & B_0 + 22B_1 + 484B_2 - 54,7 = \varepsilon_4, \\ & B_0 + 27B_1 + 729B_2 - 47,5 = \varepsilon_5, \\ 3) \quad & B_0 + 32B_1 + 1024B_2 - 41,4 = \varepsilon_6, \\ & B_0 + 37B_1 + 1369B_2 - 36,3 = \varepsilon_7. \end{aligned}$$

Отсюда тем же приёмом, как и в первом случае, получаем такие уравнения:

$$\begin{aligned} 2B_0 + 19B_1 + 193B_2 &= 156,6, \\ 3B_0 + 66B_1 + 1502B_2 &= 165,4, \\ 2B_0 + 69B_1 + 2393B_2 &= 77,7. \end{aligned}$$

Решая эти уравнения, находим:

$$\begin{aligned} B_0 &= 100,8784, \\ B_1 &= -2,6164, \\ B_2 &= 0,0236. \end{aligned}$$

Такие значения коэффициентов B_0 , B_1 и B_2 мы получаем из второй группировки уравнений.

III. Третья группировка:

$$\begin{aligned} 1) \quad & B_0 + 7B_1 + 49B_2 - 83,7 = \varepsilon_1, \\ & B_0 + 12B_1 + 144B_2 - 72,9 = \varepsilon_2, \\ & B_0 + 17B_1 + 289B_2 - 63,2 = \varepsilon_3, \\ 2) \quad & B_0 + 22B_1 + 484B_2 - 54,7 = \varepsilon_4, \\ & B_0 + 27B_1 + 729B_2 - 47,5 = \varepsilon_5, \\ 3) \quad & B_0 + 32B_1 + 1024B_2 - 41,4 = \varepsilon_6, \\ & B_0 + 37B_1 + 1369B_2 - 36,3 = \varepsilon_7. \end{aligned}$$

Отсюда получаем такие три уравнения:

$$\begin{aligned} 3B_0 + 36B_1 + 482B_2 &= 219,8, \\ 2B_0 + 49B_1 + 1213B_2 &= 102,2, \\ 2B_0 + 69B_1 + 2393B_2 &= 77,7. \end{aligned}$$

Решая эти уравнения, находим:

$$\begin{aligned} B_0 &= 100,8295, \\ B_1 &= -2,6115, \\ B_2 &= 0,0235. \end{aligned}$$

Эти значения коэффициентов B_0 , B_1 и B_2 мы получаем из третьей группировки уравнений.

Таким образом, три исследованные группировки уравнений приводят к трём эмпирическим формулам такого вида:

$$\begin{aligned} \text{I. } A &= 100,96 - 2,628t + 0,0239t^2, \\ \text{II. } A &= 100,88 - 2,616t + 0,0236t^2, \\ \text{III. } A &= 100,83 - 2,611(5)t + 0,0235t^2. \end{aligned} \quad (24)$$

Эти формулы необходимо проверить при помощи того же приёма, который мы применяли выше, т. е. подставить в их правую часть последовательно все табличные значения t и определить соответствующие значения A . Результаты такой проверки для всех трёх формул приведены в таблице XXXVII, где указаны также значения ε , которые были вычислены по формулам (23).

Из этой таблицы мы видим, что все три формулы дают очень хорошее совпадение с опытными данными, значительно лучшее, чем в таблице XXXVI. Но из этих трёх

Таблица XXXVII

А таб- личное	Формула I		Формула II		Формула III	
	А	ε	А	ε	А	ε
83,7	83,73	+0,03	83,72	+0,02	83,70	0,00
72,9	72,87	-0,03	72,89	-0,01	73,88	-0,02
63,2	63,19	-0,01	63,23	+0,03	63,23	+0,03
54,7	54,71	+0,01	54,75	+0,05	54,75	+0,05
47,5	47,43	-0,07	47,45	-0,05	47,45	-0,05
41,4	41,34	-0,06	41,33	-0,07	41,33	-0,07
36,3	36,44	+0,14	36,40	+0,10	36,38	+0,08

формул следует выбрать лучшую. Для этого принято пользоваться условием (4), т. е. определить величину $\sum \epsilon^2$, и отдавать предпочтение той формуле, для которой эта величина окажется наименьшей. Вычисляя эту величину для трёх формул (24), находим такие результаты:

$$\text{I. } \sum \epsilon^2 = 0,0301,$$

$$\text{II. } \sum \epsilon^2 = 0,0213,$$

$$\text{III. } \sum \epsilon^2 = 0,0176.$$

Мы видим, что в данном случае следует остановиться на третьей формуле как на наиболее точной:

$$A = 100,83 - 2,6115t + 0,0235t^2. \quad (24^*)$$

Если бы мы применили условие (4) к данным таблицы XXXVI и вычислили бы для них величину $\sum \epsilon^2$, то получили бы такое значение:

$$\sum \epsilon^2 = 0,0575.$$

Этот результат показывает, что формулы (24), даже наименее удачная из них, т. е. первая из формул (24), являются значительно более точными, чем формула (19).

Необходимо обратить внимание также на то обстоятельство, что если мы, применяя условие (20), вычисляем алгебраическую сумму отклонений ϵ для каждой группы

уравнений в отдельности, то получаются величины, равные нулю, или очень малые, как это видно из таблицы XXXVIII, где значения ϵ разбиты на группы соответственно группам уравнений.

Таблица XXXVIII

ε		
I	II	III
+0,03	+0,02	0,00
-0,03	-0,01	-0,02
		+0,03
-0,01	+0,03	
+0,01	+0,05	+0,05
	-0,05	-0,05
-0,07		
-0,06	-0,07	-0,07
+0,14	+0,10	+0,08

Метод средних, как видно из разобранных примера, является очень простым, не требует особенно сложных вычислений и может обеспечить хорошие результаты. Метод средних оказывается очень надёжным и даёт очень удачные формулы в особенности в тех случаях, когда число экспериментальных точек достаточно велико, так что в каждую группу уравнений входит большое число их, но во всяком случае не менее трёх-четырёх. Однако начальные уравнения можно группировать различными способами, что вызывает некоторую неопределённость результатов, так как при большом числе начальных уравнений исследовать все группировки обычно не представляется возможным. Кроме того, необходимо дополнительное исследование тех нескольких формул, к которым приводит метод средних, как это было показано на разобранных примерах. От этого недостатка свободен только метод наименьших квадратов, при котором получается только одна формула и притом лучшая из всех возможных. Но вместе

с тем метод наименьших квадратов обыкновенно требует очень сложных вычислений, которые при большом числе начальных уравнений иногда могут оказаться чрезвычайно утомительными.

5. Метод наименьших квадратов. Основное положение метода наименьших квадратов, как уже было сказано, состоит в том, что для лучшей кривой сумма квадратов отклонений ϵ (рис. 60) должна быть наименьшей:

$$\sum \epsilon^2 = \text{минимум.} \quad (4)$$

Для того чтобы метод наименьших квадратов обосновать математически, допускаем попрежнему, что мы имеем n экспериментальных точек:

$$(x_1, y_1), (x_2, y_2), (x_3, y_3), \dots, (x_n, y_n),$$

и что эмпирическая формула, которую мы ищем для этих данных, имеет вид многочлена (3).

Подставляя в эту формулу последовательно все значения x и y , мы получаем n начальных уравнений, которые имели и выше:

$$\begin{aligned} B_0 + B_1 x_1 + B_2 x_1^2 + \dots + B_m x_1^m - y_1 &= \varepsilon_1, \\ B_0 + B_1 x_2 + B_2 x_2^2 + \dots + B_m x_2^m - y_2 &= \varepsilon_2, \\ B_0 + B_1 x_3 + B_2 x_3^2 + \dots + B_m x_3^m - y_3 &= \varepsilon_3, \\ \vdots & \\ \vdots & \\ B_0 + B_1 x_n + B_2 x_n^2 + \dots + B_m x_n^m - y_n &= \varepsilon_n, \end{aligned} \quad (22)$$

причём мы считаем попрежнему, что n значительно больше, чем $(m + 1)$.

На основании этих уравнений мы видим, что условие (4), если его написать в развёрнутом виде, приводит к такому выражению:

[illegible]

Переменными величинами в этом выражении являются коэффициенты $B_0, B_1, B_2, \dots, B_m$, и мы должны найти для них такие значения, при которых выражение имеет наименьшую величину. Для этого следует воспользоваться общим приёмом дифференциального исчисления, т. е. найти частные производные выражения (25) по всем коэффициентам B и эти производные приравнять нулю. Таким образом, обозначая выражение (25) через F и вычисляя частные производные, получаем:

$$\begin{aligned}\frac{\partial F}{\partial B_0} &= 2(B_0 + B_1x_1 + B_2x_1^2 + \dots + B_mx_1^m - y_1) + \\ &\quad + 2(B_0 + B_1x_2 + B_2x_2^2 + \dots + B_mx_2^m - y_2) + \\ &\quad + 2(B_0 + B_1x_3 + B_2x_3^2 + \dots + B_mx_3^m - y_3) + \dots \\ &\quad \dots + 2(B_0 + B_1x_n + B_2x_n^2 + \dots + B_mx_n^m - y_n), \\ \frac{\partial F}{\partial B_1} &= 2(B_0 + B_1x_1 + B_2x_1^2 + \dots + B_mx_1^m - y_1) \cdot x_1 + \\ &\quad + 2(B_0 + B_1x_2 + B_2x_2^2 + \dots + B_mx_2^m - y_2) \cdot x_2 + \\ &\quad + 2(B_0 + B_1x_3 + B_2x_3^2 + \dots + B_mx_3^m - y_3) \cdot x_3 + \dots \\ &\quad \dots + 2(B_0 + B_1x_n + B_2x_n^2 + \dots + B_mx_n^m - y_n) \cdot x_n, \\ \frac{\partial F}{\partial B_2} &= 2(B_0 + B_1x_1 + B_2x_1^2 + \dots + B_mx_1^m - y_1) \cdot x_1^2 + \\ &\quad + 2(B_0 + B_1x_2 + B_2x_2^2 + \dots + B_mx_2^m - y_2) \cdot x_2^2 + \\ &\quad + 2(B_0 + B_1x_3 + B_2x_3^2 + \dots + B_mx_3^m - y_3) \cdot x_3^2 + \dots \\ &\quad \dots + 2(B_0 + B_1x_n + B_2x_n^2 + \dots + B_mx_n^m - y_n) \cdot x_n^2, \\ &\quad \vdots \\ &\quad \vdots \\ \frac{\partial F}{\partial B_m} &= 2(B_0 + B_1x_1 + B_2x_1^2 + \dots + B_mx_1^m - y_1) \cdot x_1^m + \\ &\quad + 2(B_0 + B_1x_2 + B_2x_2^2 + \dots + B_mx_2^m - y_2) \cdot x_2^m + \\ &\quad + 2(B_0 + B_1x_3 + B_2x_3^2 + \dots + B_mx_3^m - y_3) \cdot x_3^m + \dots \\ &\quad \dots + 2(B_0 + B_1x_n + B_2x_n^2 + \dots + B_mx_n^m - y_n) \cdot x_n^m.\end{aligned}$$

Приравнивая эти выражения нулю и сокращая их на два, получаем следующие уравнения:

$$\begin{aligned} & (B_0 + B_1 x_1 + B_2 x_1^2 + \dots + B_m x_1^m - y_1) + \\ & \quad + (B_0 + B_1 x_2 + B_2 x_2^2 + \dots + B_m x_2^m - y_2) + \\ & \quad + (B_0 + B_1 x_3 + B_2 x_3^2 + \dots + B_m x_3^m - y_3) + \dots \\ & \quad \dots + (B_0 + B_1 x_n + B_2 x_n^2 + \dots + B_m x_n^m - y_n) = 0, \end{aligned} \quad (26)$$

$$\begin{aligned} & (B_0 + B_1 x_1 + B_2 x_1^2 + \dots + B_m x_1^m - y_1) \cdot x_1 + \\ & \quad + (B_0 + B_1 x_2 + B_2 x_2^2 + \dots + B_m x_2^m - y_2) x_2 + \\ & \quad + (B_0 + B_1 x_3 + B_2 x_3^2 + \dots + B_m x_3^m - y_3) \cdot x_3 + \dots \\ & \quad \dots + (B_0 + B_1 x_n + B_2 x_n^2 + \dots + B_m x_n^m - y_n) \cdot x_n = 0, \\ & (B_0 + B_1 x_1 + B_2 x_1^2 + \dots + B_m x_1^m - y_1) \cdot x_1^2 + \\ & \quad + (B_0 + B_1 x_2 + B_2 x_2^2 + \dots + B_m x_2^m - y_2) x_2^2 + \\ & \quad + (B_0 + B_1 x_3 + B_2 x_3^2 + \dots + B_m x_3^m - y_3) \cdot x_3^2 + \dots \\ & \quad \dots + (B_0 + B_1 x_n + B_2 x_n^2 + \dots + B_m x_n^m - y_n) \cdot x_n^2 = 0, \\ & \dots\dots\dots \\ & (B_0 + B_1 x_1 + B_2 x_1^2 + \dots + B_m x_1^m - y_1) \cdot x_1^m + \\ & \quad + (B_0 + B_1 x_2 + B_2 x_2^2 + \dots + B_m x_2^m - y_2) x_2^m + \dots \\ & \quad \dots + (B_0 + B_1 x_n + B_2 x_n^2 + \dots + B_m x_n^m - y_n) \cdot x_n^m = 0. \end{aligned} \tag{26}$$

При помощи небольших преобразований эти уравнения можно представить также в таком виде:

[illegible]

Уравнения (26) или (26') называются нормальными уравнениями. Число нормальных урав-

нений равно, очевидно, $(m+1)$ и они содержат $(m+1)$ неизвестных коэффициентов B , по отношению к которым эти уравнения являются линейными. Решая эту систему линейных уравнений одним из способов, рекомендуемых для этого, мы находим значения всех коэффициентов B .

Рассматривая нормальные уравнения (26), нетрудно видеть, что все эти уравнения составляются по вполне определённой схеме, которую легко запомнить. Действительно, мы видим, что:

1) Для составления первого нормального уравнения следует левую часть каждого из начальных уравнений (22) умножить на коэффициент, который стоит при первом неизвестном B_0 в том же начальном уравнении, затем все полученные уравнения следует сложить и эту сумму положить равной нулю. Если коэффициенты при B_0 во всех начальных уравнениях оказываются равными единице, как это имеет место в системе начальных уравнений (22), то первое нормальное уравнение получается непосредственно в результате сложения всех начальных уравнений.

2) Точно так же для составления второго нормального уравнения следует левую часть каждого из начальных уравнений умножить на коэффициент, который стоит при втором неизвестном B_1 в том же начальном уравнении, затем все полученные в результате этого уравнения следует сложить и эту сумму приравнять нулю.

3) Совершенно тем же приёмом составляются все следующие нормальные уравнения: левые части всех начальных уравнений умножают последовательно на коэффициенты, которые стоят в этих уравнениях соответственно при B_2 (для третьего нормального уравнения), при B_3 (для четвертого нормального уравнения) и т. д. Все уравнения, получаемые в результате этих умножений, каждый раз складывают и эти суммы приравливают нулю.

Также нетрудно установить схему, по которой составляются коэффициенты при неизвестных B и свободные члены в нормальных уравнениях (26'). Действительно,

если коэффициенты при B_0 во всех начальных уравнениях равны единице, то:

1) В первом начальном уравнении (26'):

а) коэффициентом при B_0 служит число n , т. е. число начальных уравнений; б) коэффициентами при $B_1, B_2, \dots$ служит сумма всех x , причём каждое из слагаемых этой суммы входит в степени, равной индексу при B , и в) свободный член равен сумме всех y .

2) Во втором и следующих нормальных уравнениях:

а) коэффициентами при B служит попрежнему сумма всех x , но каждое из слагаемых этих сумм входит в степени: во втором нормальном уравнении на единицу больше соответствующих индексов при B , в третьем нормальном уравнении на две единицы больше индексов при B и т. д.; можно сказать, что при переходе от первого нормального уравнения к следующим происходит как бы передвижение коэффициентов при B каждый раз на один член влево;

б) свободные члены образованы попарными произведениями соответственных значений x и y , причём в этих произведениях все x входят в степенях на единицу ниже порядкового номера нормального уравнения (в правой части первого нормального уравнения все x входят в нулевой степени).

Рассмотрим два примера составления эмпирических формул методом наименьших квадратов.

Пример 1. Определим этим методом коэффициенты эмпирической формулы для той химической реакции, которую мы подробно разбирали выше. В этом случае мы ограничились в многочлене (3) тремя первыми членами и получили начальные уравнения такого вида:

$$\begin{aligned} B_0 + 7B_1 + 49B_2 - 83,7 &= \varepsilon_1, \\ B_0 + 12B_1 + 144B_2 - 72,9 &= \varepsilon_2, \\ B_0 + 17B_1 + 289B_2 - 63,2 &= \varepsilon_3, \\ B_0 + 22B_1 + 484B_2 - 54,7 &= \varepsilon_4, \\ B_0 + 27B_1 + 729B_2 - 47,5 &= \varepsilon_5, \\ B_0 + 32B_1 + 1024B_2 - 41,4 &= \varepsilon_6, \\ B_0 + 37B_1 + 1369B_2 - 36,3 &= \varepsilon_7. \end{aligned}$$

Применяя уравнения (26') и те указания, которые по отношению к ним были сделаны, находим выражения, определяющие нормальные уравнения:

$$\begin{aligned} 1) \quad & 7B_0 + (7 + 12 + 17 + 22 + 27 + 32 + 37)B_1 + \\ & + (7^2 + 12^2 + 17^2 + 22^2 + 27^2 + 32^2 + 37^2)B_2 = \\ & = 83,7 + 72,9 + 63,2 + 54,7 + 47,5 + 41,4 + 36,3; \\ 2) \quad & (7 + 12 + 17 + 22 + 27 + 32 + 37)B_0 + \\ & + (7^2 + 12^2 + 17^2 + 22^2 + 27^2 + 32^2 + 37^2)B_1 + \\ & + (7^3 + 12^3 + 17^3 + 22^3 + 27^3 + 32^3 + 37^3)B_2 = \\ & = 83,7 \cdot 7 + 72,9 \cdot 12 + 63,2 \cdot 17 + 54,7 \cdot 22 + \\ & + 47,5 \cdot 27 + 41,4 \cdot 32 + 36,3 \cdot 37; \\ 3) \quad & (7^2 + 12^2 + 17^2 + 22^2 + 27^2 + 32^2 + 37^2)B_0 + \\ & + (7^3 + 12^3 + 17^3 + 22^3 + 27^3 + 32^3 + 37^3)B_1 + \\ & + (7^4 + 12^4 + 17^4 + 22^4 + 27^4 + 32^4 + 37^4)B_2 = \\ & = 83,7 \cdot 7^2 + 72,9 \cdot 12^2 + 63,2 \cdot 17^2 + 54,7 \cdot 22^2 + \\ & + 47,5 \cdot 27^2 + 41,4 \cdot 32^2 + 36,3 \cdot 37^2. \end{aligned}$$

Выполняя все указанные действия, получаем нормальные уравнения такого вида:

$$\begin{aligned} 1) \quad & 7B_0 + 154B_1 + 4088B_2 = 399,7, \\ 2) \quad & 154B_0 + 4088B_1 + 120\,736B_2 = 7688,9, \\ 3) \quad & 4088B_0 + 120\,736B_1 + 3\,795\,092B_2 = 186\,054,3. \end{aligned}$$

Решая эти уравнения, получаем такие значения коэффициентов B_0, B_1 и B_2 :

$$\begin{aligned} B_0 &= 100,791, \\ B_1 &= -2,6066, \\ B_2 &= 0,023381. \end{aligned}$$

Таким образом, эмпирическая формула, рассчитанная методом наименьших квадратов, имеет такой вид:

$$A = 100,79 - 2,6066t + 0,02338t^2. \quad (27)$$

В этой формуле значения коэффициентов B несколько отличаются от их значений, которые мы получали раньше, применяя другие методы, например метод средних. Вместе

с тем мы видим, что значения коэффициентов в формуле (27) ближе всего отвечают их значениям в третьей из формул (24), которая из прежних формул оказалась наиболее точной. Для того чтобы сравнить формулы (24) с формулой (27), следует подставить в последнюю все табличные значения t и вычислить соответствующие значения A , а также значения отклонений ϵ . Результаты этих вычислений приведены в таблице XXXIX.

Таблица XXXIX

A табличное	A вычисленное	ϵ
83,7	83,69	-0,01
72,9	72,88	-0,02
63,2	63,24	+0,04
54,7	54,76	+0,06
47,5	47,46	-0,04
41,4	41,32	-0,08
36,3	36,36	+0,06

Из этой таблицы мы видим, что формула (27) даёт очень хорошее совпадение с экспериментальными значениями A . Вместе с тем, вычисляя величину $\sum \epsilon^2$, мы получаем:

$$\sum \epsilon^2 = 0,0173,$$

т. е. величину, ещё несколько меньшую той, которую мы нашли для третьей из формул (24). Таким образом, мы имеем все основания считать, что формула (27) является лучшей формулой из всех многочисленных формул, полученных нами для данной химической реакции.

Пример 2. На графике нанесён ряд точек (рис. 61), координаты которых даны в таблице XL.

Таблица XL

x	0,5	1,0	1,5	2,0	2,5	3,0
y	0,62	1,64	2,58	3,70	5,02	6,04

Составим, применяя метод наименьших квадратов, уравнение прямой, которая проходит возможно ближе ко всем этим точкам.

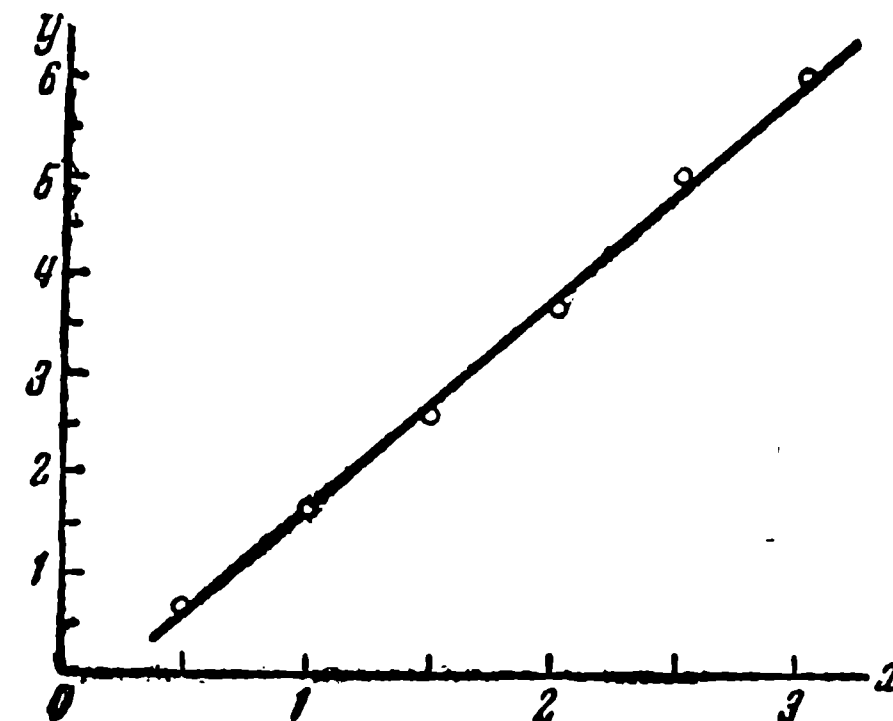


Рис. 61.

Уравнение этой прямой мы можем написать в таком виде:

$$y = B_0 + B_1 x.$$

Подставляя в это выражение значения x и y , находим начальные уравнения:

$$B_0 + 0,5B_1 - 0,62 = \epsilon_1,$$

$$B_0 + 1,0B_1 - 1,64 = \epsilon_2,$$

$$B_0 + 1,5B_1 - 2,58 = \epsilon_3,$$

$$B_0 + 2,0B_1 - 3,70 = \epsilon_4,$$

$$B_0 + 2,5B_1 - 5,02 = \epsilon_5,$$

$$B_0 + 3,0B_1 - 6,04 = \epsilon_6.$$

Применяя выражения (26'), составляем первое и второе нормальные уравнения:

$$\begin{aligned} 6B_0 + (0,5 + 1,0 + 1,5 + 2,0 + 2,5 + 3,0) B_1 &= \\ &= 0,62 + 1,64 + 2,58 + 3,70 + 5,02 + 6,04, \\ (0,5 + 1,0 + 1,5 + 2,0 + 2,5 + 3,0) B_0 + \\ &+ (0,5^2 + 1,0^2 + 1,5^2 + 2,0^2 + 2,5^2 + 3,0^2) B_1 = \\ &= 0,62 \cdot 0,5 + 1,64 \cdot 1 + 2,58 \cdot 1,5 + 3,70 \cdot 2,0 + \\ &+ 5,02 \cdot 2,5 + 6,04 \cdot 3,0 \end{aligned}$$

или

$$\begin{aligned} 6 B_0 + 10,5 B_1 &= 19,60, \\ 10,5 B_0 + 22,75 B_1 &= 43,89. \end{aligned}$$

Решая эти уравнения, находим:

$$\begin{aligned} B_0 &= -0,57, \\ B_1 &= 2,192. \end{aligned}$$

Таким образом, уравнение искомой прямой имеет такой вид:

$$y = -0,57 + 2,192 x.$$

Эта прямая дана на рис. 61.

Метод наименьших квадратов является самым точным и может обеспечить вполне надёжные результаты, если все вычисления производить с достаточной точностью. Здесь необходимо сделать одно весьма существенное замечание относительно той точности, с которой следует вести вычисления коэффициентов в эмпирических формулах. В этих формулах имеются члены, которые содержат вторые, а иногда и более высокие степени x . Вследствие этого часто приходится коэффициенты в таких членах умножать на большие числа. Так, например, значение t в формуле (27) может достигать 37 сек., квадрат этой величины превышает 1000, и на такие большие числа мы умножаем коэффициент B_2 . В результате этого цифры множимого, которые стояли на третьем месте после запятой, в произведении переходят в разряд целых. Хотя ряд (3) является сходящимся и вследствие этого абсолютная величина коэффициентов B быстро уменьшается, тем не менее их значения необходимо давать с достаточно большим числом десятичных знаков, как это дано, например, в формуле (27). Вследствие этого рекомендуется все исходные данные, которые служат для вычисления коэффициентов B , рассматривать как числа точные, т. е. при всех вычислениях следует крайне осторожно производить

округление чисел или отбрасывать знаки, считая их лишними. Только закончив вычисления всех коэффициентов B , следует в их окончательных значениях отбросить излишние знаки. Все коэффициенты B необходимо давать с таким числом десятичных знаков, чтобы значения y , вычисляемые по формуле, имели ту же точность, как и его табличные значения. Отсюда следует, что значения коэффициентов B следует давать с различным числом десятичных знаков, которое должно возрастать с возрастанием индекса B . Вследствие этого в формуле (27) коэффициенты B_0 , B_1 и B_2 даны соответственно с двумя, четырьмя и пятью десятичными знаками.

6. Выбор формул лучшего типа. Те методы определения постоянных в эмпирических формулах, которые мы разобрали, можно применять только в тех случаях, когда эти формулы представляют собой несколько первых членов разложения функции $f(x)$ в ряд по степеням x , т. е. когда эмпирическая формула имеет вид многочлена (3).

Хотя этот вид эмпирических формул, как уже было сказано, применяется чрезвычайно часто, однако встречаются и такие случаи, когда возможность применять такие формулы вызывает затруднения. Действительно, следует помнить, что, применяя ряд (3), нам необходимо, прежде всего, определить степень x , или число членов ряда, а для этого приходится составлять таблицу разностей. Это возможно только в том случае, когда значения x образуют арифметическую прогрессию (стр. 261). и если последнее не имеет места, то необходимо вычислять значения y для промежуточных значений x , применяя интерполирование при неравноотстоящих значениях аргумента (стр. 285), что может внести осложнения и понизить точность результатов. Вследствие этого часто приходится находить и другие виды алгебраических выражений, которые могли бы с достаточной точностью представлять данный ряд измерений. Примеры таких эмпирических формул мы уже встречали (стр. 337).

Не имеется общего метода, который давал бы возможность находить эмпирические формулы лучшего типа

для любого ряда измерений. Но вместе с тем можно дать ряд общих указаний, которые очень облегчают выбор эмпирических формул лучшего типа, если применение формул вида (3) в силу каких-либо причин оказывается невозможным. Эти указания в основном сводятся к следующему:

1) Необходимо иметь на обычной прямоугольной координатной сетке графические изображения различных кривых, которые ниже указываются, вычерченными в одном и том же масштабе. Экспериментальный материал, для которого мы подыскиваем эмпирическую формулу, следует, также изобразить графически на прямоугольной координатной сетке в том же масштабе. При этих условиях непосредственное сравнение экспериментального графика с образцами кривых обыкновенно даёт достаточно определённые указания на типы эмпирических зависимостей, возможных для данного ряда измерений. После этого можно приступать к предварительным вычислениям.

2) Если экспериментальные точки, нанесённые на прямоугольной координатной сетке, располагаются на прямой линии или очень близко ей соответствуют, то следует остановиться на линейной зависимости, т. е. положить:

$$y = a + bx,$$

и определить затем параметры прямой a и b одним из указанных выше способов.

3) Если экспериментальные точки на графике не располагаются на прямой, то следует перенести их на логарифмическую и на полулогарифмическую сетки. Очень часто оказывается, что на одной из этих сеток точки ложатся на прямую или очень близко к ней. В таком случае:

а) Если точки образуют прямую на логарифмической сетке, то следует применить формулу

$$y = ax^n + b,$$

где в некоторых случаях b может оказаться равным нулю.

б) Если точки образуют прямую на полулогарифмической сетке, то следует применить формулу

$$y = ae^{nx} + b,$$

где b иногда может оказаться также равным нулю.

Примеры таких эмпирических формул мы уже имели и познакомились с методами определения их постоянных (стр. 337 и след.).

4) Если экспериментальные точки не располагаются на прямой ни на одной из указанных координатных сеток, то необходимо переходить к другим видам алгебраических формул. В общем случае они могут быть весьма разнообразными, но если на прямоугольной координатной сетке экспериментальная кривая является вполне плавной без резких искривлений, то может оказаться подходящей одна из следующих формул.

а) Формула вида

$$y = a + \frac{b}{x}. \quad (28)$$

В этом случае, полагая

$$x = \frac{1}{z},$$

мы получаем:

$$y = a + bz,$$

т. е. формулу, которая отвечает прямой линии. Иначе говоря, если мы вычислим значения z и нанесём точки (y, z) в прямоугольной системе координат, то должна получиться прямая. Если этого не наблюдается, то формулу (28) следует считать для данного ряда измерений непригодной.

б) Формула вида

$$y = \frac{1}{a + bx}. \quad (29)$$

В этом случае, полагая

$$y = \frac{1}{z},$$

находим:

$$z = a + bx,$$

т. е. опять прямую линию, что даёт возможность непосредственно определить пригодность формулы (29).

в) Формула вида

$$\lg y = a + bx + cx^2. \quad (30)$$

Эта формула является линейной относительно постоянных a , b и c , и для их определения можно воспользоваться одним из указанных выше способов.

г) Формула вида

$$y = \frac{1}{a + bx + cx^2}. \quad (31)$$

Полагая здесь

$$y = \frac{1}{z},$$

находим:

$$z = a + bx + cx^2,$$

т. е. формулу вида (3).

д) Формула вида

$$y = a + \frac{b}{x} + \frac{c}{x^2}. \quad (32)$$

Полагая здесь

$$\frac{1}{x} = z,$$

находим:

$$y = a + bz + cz^2,$$

т. е. вновь получается формула вида (3).

е) Формула вида

$$y = \frac{x}{a + bx + cx^2}. \quad (33)$$

Полагая здесь

$$\frac{x}{y} = z,$$

получаем после небольших преобразований:

$$z = a + bx + cx^2,$$

т. е. формулу вида (3).

ж) Формула вида

$$y = ae^{nx+mx^2}. \quad (34)$$

Логарифмируя, находим:

$$\lg y = \lg a + nx \lg e + mx^2 \lg e;$$

полагая здесь

$$\lg y = z, \quad \lg a = p, \quad n \lg e = q \quad \text{и} \quad m \lg e = r,$$

находим:

$$z = p + qx + rx^2,$$

т. е. формулу вида (3).

Рекомендуется также много других формул, но на них мы останавливаться не будем.

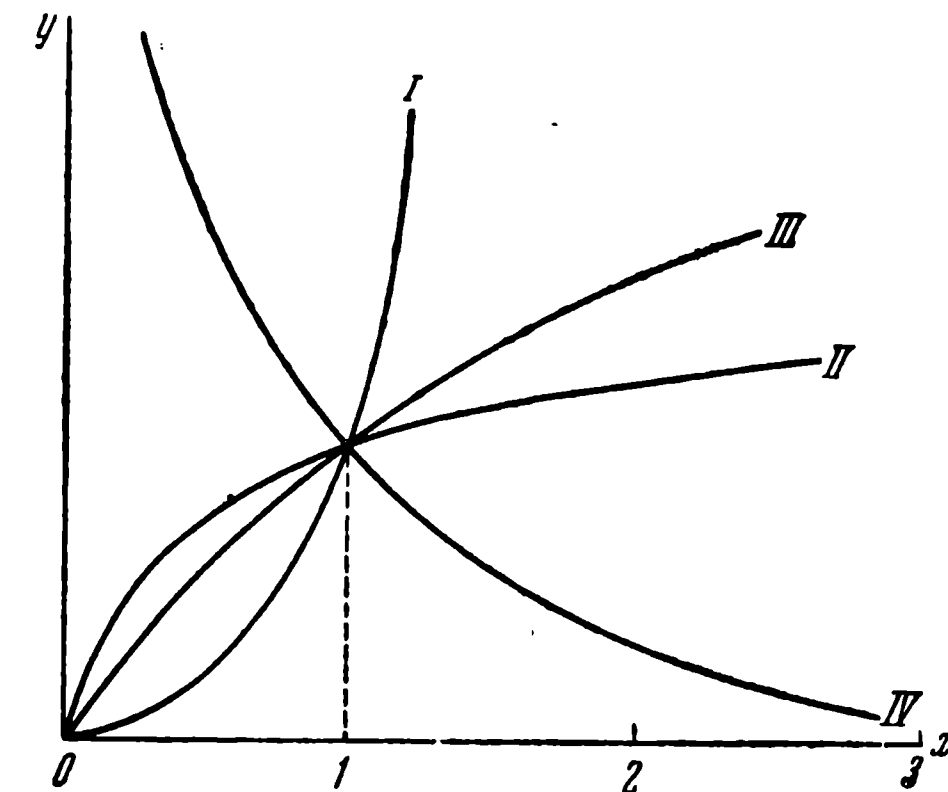


Рис. 62.

Пригодность формул (28) — (34) для данного ряда измерений можно непосредственно проверить в тех случаях, когда график кривой может быть преобразован в прямую линию. Кроме того, пригодность формул можно определить сравнением соответствующих кривых с экспериментальным графиком. Здесь следует иметь в виду, что вид кривой может очень сильно изменяться при изменении числовых значений параметров в её уравнении. Так, на

рис. 62 дано графическое изображение степенной функции

$$y = ax^n$$

при различных значениях n . При n положительном и большем единицы мы имеем параболическую кривую (кривая *I*), при n положительном, но меньшем единицы, кривая может иметь вид *II* и *III*, при n отрицательном получаются кривые гиперболического вида (кривая *IV*). Таким образом, указанные кривые дают очень широкие возможности выбора эмпирической формулы лучшего типа, но удачное решение этих вопросов, т. е. выбор лучшей формулы и определение её коэффициентов, во многом зависит от практики и опытности в таких вычислительных работах.

ПРИЛОЖЕНИЯ

1. Формулы для приближённых вычислений

Если $\alpha, \beta, \gamma, \delta, \dots$ малы сравнительно с единицей, то

$$\begin{aligned} (1 \pm \alpha)(1 \pm \beta) &= 1 \pm \alpha \pm \beta, \\ (1 \pm \alpha)(1 \pm \beta)(1 \pm \gamma)(1 \pm \delta) \dots &= 1 \pm \alpha \pm \beta \pm \gamma \pm \delta \pm \dots, \\ (1 \pm \alpha)^2 &= 1 \pm 2\alpha, & e^\alpha &= 1 + \alpha, \\ (1 \pm \alpha)^n &= 1 \pm n\alpha, & a^\alpha &= 1 + \alpha \ln a, \\ \sqrt{1 \pm \alpha} &= 1 \pm \frac{\alpha}{2}, & \sqrt[n]{1 \pm n\alpha} &= 1 \pm \alpha, \\ \frac{1}{\sqrt{1 \pm \alpha}} &= 1 \mp \frac{\alpha}{2}, & \frac{1}{1 \pm \alpha} &= 1 \mp \alpha, \\ \frac{(1 \pm \alpha)(1 \pm \beta) \dots}{(1 \pm \gamma)(1 \pm \delta) \dots} &= 1 \pm \alpha \pm \beta \mp \dots \mp \gamma \mp \delta \dots \end{aligned}$$

Если a и b близки между собой, то

$$\sqrt{ab} = \frac{1}{2}(a + b).$$

Если α выражено в долях радиана и мало по сравнению с ним, то

$$\begin{aligned} \sin \alpha &= \operatorname{tg} \alpha = \alpha, \\ \cos \alpha &= 1, \\ \sin(\varphi \pm \alpha) &= \sin \varphi \pm \alpha \cos \varphi, \\ \cos(\varphi \mp \alpha) &= \cos \varphi \mp \alpha \sin \varphi. \end{aligned}$$

2. Абсолютные и относительные ошибки некоторых функций

Функция y определяется выражением

$$y = f(x_1, x_2, x_3, \dots).$$

Абсолютные ошибки величин $y, x_1, x_2, x_3, \dots$ обозначены соответственно через

$$dy, dx_1, dx_2, dx_3, \dots;$$

их относительные ошибки δ обозначены через

$$\frac{dy}{y}, \frac{dx_1}{x_1}, \frac{dx_2}{x_2}, \frac{dx_3}{x_3}, \dots$$

Вид функции y	Абсолютная ошибка dy	Относительная ошибка $\frac{dy}{y}$
$y = x_1 + x_2 + x_3 + \dots$	$\pm (dx_1 + dx_2 + dx_3 + \dots)$	$\pm \frac{dx_1 + dx_2 + dx_3 + \dots}{x_1 + x_2 + x_3 + \dots}$
$y = x_1 - x_2$	$\pm (dx_1 + dx_2)$	$\pm \frac{dx_1 + dx_2}{x_1 - x_2}$
$y = x_1 \cdot x_2$	$\pm (x_1 dx_2 + x_2 dx_1)$	$\pm \left(\frac{dx_1}{x_1} + \frac{dx_2}{x_2} \right)$
$y = x_1 \cdot x_2 \cdot x_3$	$\pm (x_2 x_3 dx_1 + x_3 x_1 dx_2 + x_1 x_2 dx_3)$	$\pm \left(\frac{dx_1}{x_1} + \frac{dx_2}{x_2} + \frac{dx_3}{x_3} \right)$
$y = ax$	$\pm a dx$	$\pm \frac{dx}{x}$
$y = x^n$	$\pm nx^{n-1} \cdot dx$	$\pm n \frac{dx}{x}$
$y = \sqrt[n]{x}$	$\pm \frac{1}{n} x^{\frac{1}{n}-1} \cdot dx$	$\pm \frac{1}{n} \frac{dx}{x}$
$y = \sin x$	$\pm \cos x dx$	$\pm \operatorname{ctg} x dx$
$y = \cos x$	$\pm \sin x dx$	$\pm \operatorname{tg} x dx$
$y = \operatorname{tg} x$	$\pm \frac{dx}{\cos^2 x}$	$\pm \frac{2dx}{\sin 2x}$
$y = \operatorname{ctg} x$	$\pm \frac{dx}{\sin^2 x}$	$\pm \frac{2dx}{\sin 2x}$

3. Значения интеграла вероятностей $P = \frac{2}{\sqrt{\pi}} \int_0^z e^{-x^2} dx$, где $z = hx$, в пределах от $hx=0,00$ до $hx=0,99$

$z = hx$	0	1	2	3	4	5	6	7	8	9
0,00	0,00000	00 113	00 226	00 339	00 451	00 564	00 677	00 790	00 903	01 016
0,01	0,01128	01 241	01 354	01 467	01 580	01 792	01 805	01 918	02 031	02 144
0,02	0,02256	02 369	02 482	02 595	02 708	02 820	02 933	03 046	03 159	03 271
0,03	0,03384	03 497	03 610	03 722	03 835	03 948	04 060	04 173	04 286	04 398
0,04	0,04511	04 624	04 736	04 849	04 962	05 074	05 187	05 299	05 412	05 525
0,05	0,05637	05 750	05 862	05 975	06 087	06 200	06 312	06 425	06 537	06 650
0,06	0,06762	06 875	06 987	07 099	07 212	07 324	07 437	07 549	07 661	07 773
0,07	0,07886	07 998	08 110	08 223	08 335	08 447	08 559	08 671	08 784	08 896
0,08	0,09008	09 120	09 232	09 344	09 456	09 568	09 680	09 792	09 904	10 016
0,09	0,10128	10 240	10 352	10 464	10 576	10 687	10 799	10 911	11 023	11 135
0,10	0,11246	11 358	11 470	11 581	11 693	11 805	11 916	12 028	12 139	12 251
0,11	0,12362	12 474	12 585	12 697	12 808	12 915	13 031	13 142	13 253	13 365
0,12	0,13476	13 587	13 698	13 809	13 921	14 032	14 143	14 254	14 365	14 476
0,13	0,14587	14 698	14 809	14 919	15 030	15 141	15 252	15 363	15 473	15 584
0,14	0,15695	15 805	15 916	16 027	16 137	16 248	16 358	16 468	16 579	16 689
0,15	0,16800	16 910	17 020	17 130	17 241	17 351	17 461	17 571	17 681	17 791
0,16	0,17901	18 011	18 121	18 231	18 341	18 451	18 560	18 670	18 780	18 890
0,17	0,18999	19 109	19 218	19 328	19 437	19 547	19 656	19 766	19 875	19 984
0,18	0,20094	20 203	20 312	20 421	20 530	20 639	20 748	20 857	20 966	21 075
0,19	0,21184	21 293	21 402	21 510	21 619	21 728	21 836	21 945	22 053	22 162
0,20	0,22270	22 379	22 487	22 595	22 704	22 812	22 920	23 028	23 136	23 244
0,21	0,23352	23 460	23 568	23 676	23 784	23 891	23 999	24 107	24 214	24 322
0,22	0,24430	24 537	24 645	24 752	24 859	24 967	25 074	25 181	25 288	25 395
0,23	0,25502	25 609	25 716	25 823	25 930	26 037	26 144	26 250	26 357	26 463
0,24	0,26570	26 677	26 783	26 889	26 996	27 102	27 208	27 314	27 421	27 527

Продолжение

$z=hx$	0	1	2	3	4	5	6	7	8	9
0,25	0,27633	27 739	27 845	27 950	28 056	28 162	28 268	28 373	28 479	28 584
0,26	0,28690	28 795	28 901	29 006	29 111	29 217	29 322	29 427	29 532	29 637
0,27	0,29742	29 847	29 952	30 056	30 161	30 266	30 370	30 475	30 579	30 684
0,28	0,30788	30 892	30 997	31 101	31 205	31 309	31 413	31 517	31 621	31 725
0,29	0,31828	31 922	32 036	32 139	32 243	32 346	32 450	32 553	32 656	32 760
0,30	0,32863	32 966	33 069	33 172	33 275	33 378	33 480	33 583	33 686	33 788
0,31	0,33894	33 993	34 096	34 198	34 300	34 403	34 505	34 607	34 709	34 811
0,32	0,34913	35 014	35 116	35 218	35 319	35 421	35 523	35 624	35 725	35 827
0,33	0,35928	36 029	36 130	36 231	36 332	36 433	36 534	36 635	36 735	36 836
0,34	0,36936	37 037	37 137	37 238	37 338	37 438	37 538	37 638	37 738	37 838
0,35	0,37938	38 038	38 138	38 237	38 337	38 436	38 536	38 635	38 735	38 834
0,36	0,38933	39 032	39 131	39 230	39 329	39 428	39 526	39 625	39 724	39 822
0,37	0,39921	40 019	40 117	40 215	40 314	40 412	40 510	40 608	40 705	40 803
0,38	0,40901	40 999	41 096	41 194	41 291	41 388	41 486	41 583	41 680	41 777
0,39	0,41874	41 971	42 068	42 164	42 261	42 358	42 454	42 550	42 647	42 743
0,40	0,42839	42 935	43 031	43 127	43 223	43 319	43 415	43 510	43 606	43 701
0,41	0,43797	43 892	43 988	44 083	44 178	44 273	44 368	44 463	44 557	44 652
0,42	0,44747	44 841	44 936	45 030	45 124	45 219	45 313	45 407	45 501	45 595
0,43	0,45689	45 782	45 876	45 970	46 063	46 157	46 250	46 343	46 436	46 529
0,44	0,46623	46 715	46 808	46 901	46 994	47 086	47 179	47 271	47 364	47 456
0,45	0,47548	47 640	47 732	47 824	47 916	48 008	48 100	48 191	48 283	48 374
0,46	0,48466	48 557	48 648	48 739	48 830	48 921	49 012	49 103	49 193	49 284
0,47	0,49375	49 465	49 555	49 646	49 736	49 826	49 916	50 006	50 096	50 185
0,48	0,50275	50 365	50 454	50 543	50 633	50 722	50 811	50 900	50 989	51 078
0,49	0,51167	51 256	51 344	51 433	51 521	51 609	51 698	51 786	51 874	51 962

Продолжение

$z=hx$	0	1	2	3	4	5	6	7	8	9
0,50	0,52050	52 138	52 226	52 313	52 401	52 488	52 576	52 663	52 750	52 837
0,51	0,52924	53 011	53 098	53 185	53 272	53 358	53 445	53 531	53 617	53 704
0,52	0,53790	53 876	53 962	54 048	54 134	54 219	54 305	54 390	54 476	54 561
0,53	0,54646	54 732	54 817	54 902	54 987	55 071	55 156	55 241	55 325	55 410
0,54	0,55494	55 578	55 662	55 746	55 830	55 914	55 998	56 082	56 165	56 249
0,55	0,56332	56 416	56 499	56 582	56 665	56 748	56 831	56 914	56 996	57 079
0,56	0,57162	57 244	57 326	57 409	57 491	57 573	57 655	57 737	57 818	57 900
0,57	0,57982	58 063	58 144	58 226	58 307	58 388	58 469	58 550	58 631	58 712
0,58	0,58792	58 873	58 953	59 034	59 114	59 194	59 274	59 354	59 434	59 514
0,59	0,59594	59 673	59 753	59 832	59 912	59 991	60 070	60 149	60 228	60 307
0,60	0,60386	60 464	60 543	60 621	60 700	60 778	60 856	60 934	61 012	61 090
0,61	0,61168	61 246	61 323	61 401	61 478	61 556	61 633	61 710	61 787	61 864
0,62	0,61941	62 018	62 095	62 171	62 248	62 324	62 400	62 477	62 553	62 629
0,63	0,62705	62 780	62 856	62 932	63 007	63 083	63 158	63 233	63 309	63 384
0,64	0,63459	63 533	63 608	63 683	63 757	63 832	63 906	63 981	64 055	64 129
0,65	0,64203	64 277	64 351	64 424	64 498	64 572	64 645	64 718	64 791	64 865
0,66	0,64938	65 011	65 083	65 156	65 229	65 301	65 374	65 446	65 519	65 591
0,67	0,65663	65 735	65 807	65 878	65 950	66 022	66 093	66 165	66 236	66 307
0,68	0,66378	66 449	66 520	66 591	66 662	66 732	66 803	66 873	66 944	67 014
0,69	0,67084	67 154	67 224	67 294	67 364	67 433	67 503	67 572	67 642	67 711
0,70	0,67780	67 849	67 918	67 987	68 056	68 125	68 193	68 262	68 330	68 398
0,71	0,68467	68 535	68 603	68 671	68 738	68 806	68 874	68 941	69 009	69 076
0,72	0,69143	69 210	69 278	69 344	69 411	69 478	69 545	69 611	69 678	69 744
0,73	0,69810	69 877	69 943	70 009	70 075	70 140	70 206	70 272	70 337	70 403
0,74	0,70468	70 533	70 598	70 663	70 728	70 793	70 858	70 922	70 987	71 051

Продолжение

$z=hx$	0	1	2	3	4	5	6	7	8	9
0,75	0,71116	71 180	71 244	71 308	71 372	71 436	71 500	71 563	71 627	71 690
0,76	0,71754	71 817	71 880	71 943	72 006	72 069	72 132	72 195	72 257	72 320
0,77	0,72382	72 444	72 507	72 569	72 631	72 693	72 755	72 816	72 878	72 940
0,78	0,73001	73 062	73 124	73 185	73 246	73 307	73 368	73 429	73 489	73 550
0,79	0,73610	73 671	73 731	73 791	73 851	73 911	73 971	74 031	74 091	74 151
0,80	0,74210	74 270	74 329	74 388	74 447	74 506	74 565	74 624	74 683	74 742
0,81	0,74800	74 859	74 917	74 976	75 034	75 092	75 150	75 208	75 266	75 323
0,82	0,75381	75 439	75 496	75 553	75 611	75 668	75 725	75 782	75 839	75 896
0,83	0,75952	76 009	76 066	76 122	76 178	76 234	76 291	76 347	76 403	76 459
0,84	0,76514	76 570	76 626	76 681	76 736	76 792	76 847	76 902	76 957	77 012
0,85	0,77067	77 122	77 176	77 231	77 285	77 340	77 394	77 448	77 502	77 556
0,86	0,77610	77 664	77 718	77 771	77 825	77 878	77 932	77 985	78 038	78 091
0,87	0,78144	78 197	78 250	78 302	78 355	78 408	78 460	78 512	78 565	78 617
0,88	0,78669	78 721	78 773	78 824	78 876	78 928	78 979	79 031	79 082	79 133
0,89	0,79184	79 235	79 286	79 337	79 388	79 439	79 489	79 540	79 590	79 641
0,90	0,79691	79 741	79 791	79 841	79 891	79 941	79 990	80 040	80 090	80 139
0,91	0,80188	80 238	80 287	80 336	80 385	80 434	80 482	80 531	80 580	80 628
0,92	0,80677	80 725	80 773	80 822	80 870	80 918	80 966	81 013	81 061	81 109
0,93	0,81156	81 204	81 251	81 299	81 346	81 393	81 440	81 487	81 534	81 580
0,94	0,81627	81 674	81 720	81 767	81 813	81 859	81 905	81 951	81 997	82 043
0,95	0,82089	82 135	82 180	82 226	82 271	82 317	82 362	82 407	82 452	82 497
0,96	0,82542	82 587	82 632	82 677	82 721	82 766	82 810	82 855	82 899	82 943
0,97	0,82987	83 031	83 075	83 119	83 162	83 206	83 250	83 293	83 337	83 380
0,98	0,83423	83 466	83 509	83 552	83 595	83 638	83 681	83 723	83 766	83 808
0,99	0,83851	83 893	83 935	83 977	84 020	84 061	84 103	84 145	84 187	84 229

РЕКОМЕНДУЕМАЯ ЛИТЕРАТУРА

1. А. Н. Крылов, Лекции о приближённых вычислениях, Гостехиздат, Москва—Ленинград, 1950.
2. Дж. Скарборо, Численные методы математического анализа, Москва—Ленинград, 1934.
3. Г. Занден, Элементы прикладного анализа, Москва—Ленинград, 1932.
4. Э. Уиттекер и Г. Робинсон, Математическая обработка результатов наблюдений, Москва—Ленинград, 1933.
5. Н. А. Глаголев, Курс номографии, Гостехиздат, Москва—Ленинград, 1943.
6. Н. А. Глаголев, Учебный атлас по номографии, Москва, 1933.

ПРЕДМЕТНЫЙ УКАЗАТЕЛЬ

- Абсолютная ошибка 21
 — — корня 75
 — — произведения 68
 — — разности 63
 — — степени 73
 — — функции двух независи-
 мых переменных 88
 — — — нескольких независи-
 мых переменных 98
 — — — одного независимого
 переменного 86
 — — частного 70
 Аксиомы теории случайных оши-
 бок 143
 Анализ гармонический 294
 Анаморфоза логарифмическая
 244

 Верные цифры в приближённых
 числах 33
 Вероятная ошибка 175
 — —, геометрическое значе-
 ние 182
 Вероятное событие 131
 Вероятность, математическое
 определение 131
 — события 132
 — — невозможного 132
 — — сложного 136

 Вес измерений 192
 Возведение в степень прибли-
 жённых чисел 72
 Вторая задача общей теории
 ошибок 108
 Выбор формулы лучшего вида 367
 Вычитание приближённых чи-
 сел 63

 Гармонический анализ 294
 — —, табличный 321
 — ряд колебаний 296
 Геометрическое значение ве-
 роятной ошибки 182
 — — средней квадратичной
 ошибки 182
 — — — ошибки 184
 Графический метод 331
 Графическое дифференцирование
 212, 213
 — изображение результатов из-
 мерений 201
 — интегрирование 212, 216
 — интерполирование 259

 Деление приближённых чисел 69
 Дифференцирование графиче-
 ское 213
 Достоверное событие 130

 Единственно возможное собы-
 тие 129

 Закон нормального распределе-
 ния случайных ошибок 143,
 149
 Запись приближённых чисел 35
 Зет-номограммы 250
 Знак ошибок приближённых чи-
 сел 26
 Значения интеграла вероятно-
 стей 155, 375
 — приближённых чисел с из-
 бытком 21
 — — — с недостатком 21

 Интеграл вероятностей 149
 — —, значение 155, 375
 Интегрирование графическое
 216
 Интерполирование 257
 — графическое 259
 — линейное 259
 — при неравноотстоящих зна-
 чениях аргумента 285
 — — равноотстоящих значе-
 ниях аргумента 260
 Интерполяционных формул ме-
 тод 342

 Классификация приближённых
 чисел 29
 Координатные сетки 231
 Криволинейные номограммы 239

 Линейное интерполирование 259
 Логарифмическая анаморфоза
 244
 — шкала 226
 Лучевые номограммы 241

 Малые величины различных по-
 рядков 47
 Математическое определение ве-
 роятности 131
 Мера точности 149
 — — отклонений от среднего
 арифметического 162, 166
 — — ошибок 162
 Метод графический 331
 — интерполяционных формул
 342
 — наименьших квадратов 358
 — равных влияний 109
 — средних 349
 Механические приёмы гармониче-
 ского анализа 325
 Модуль шкалы 220

 Наиболее вероятное значение
 измеряемой величины 162
 — выгодные условия измере-
 ния 85, 120
 Наименьших квадратов метод
 358
 Невозможное событие 131
 Несовместимые события 129
 Номограммы 234
 — лучевые 241
 — на выравненных точках 245
 — сетчатые 236
 Номография 234
 Нормальное распределение слу-
 чайных ошибок 128, 149

 Обратная задача теории оши-
 бок 85, 108
 Общая классификация случай-
 ных явлений 129
 — теория ошибок функции
 80

- Округление приближённых чисел 19
- Определение наиболее выгодных условий измерения 120
- Основная формула теории случайных ошибок 149
- Основные задачи общей теории ошибок 80, 84
- Относительная ошибка 22
- — корня 74
- — произведения 67
- — разности 64
- — степени 72
- — суммы 59
- — функции двух независимых переменных 88
- — — нескольких независимых переменных 99
- — — одного независимого переменного 87
- — частного 70
- Оценка измерительных методов 111
- Ошибка абсолютная 21, 26
- вероятная 175
- квадратичная 172
- относительная 22, 27
- средняя 177
- Ошибки корня 74
- предельные 38
- произведения 66
- разности 63
- систематические 32
- случайные 32, 128
- среднего арифметического 184
- степени 72
- суммы 58
- функции 80
- частного 69
- Первая задача общей теории ошибок 84
- Периодические процессы 294
- Показатели точности измерений 159
- Порядок малых величин 47
- Правило дополнения 19
- Предельная ошибка 38
- Приближённые значения 18
- числа и величины 18
- Приближённых чисел округление 19
- Приведённые разности 285
- Проективная шкала 229
- Процентная ошибка 23
- Равномерная шкала 219
- Разности различных порядков 262
- приведённые 285
- Сетчатые номограммы 236
- Систематические ошибки 32
- Сложение приближённых чисел 57
- Сложное событие 136
- Случайные ошибки 32, 128
- явления 128
- События вероятные 131
- достоверные 130
- единственно возможные 129
- невозможные 131
- практически достоверные 133
- — невозможные 133
- Спрямление кривых 232
- Среднее арифметическое 159, 162
- —, ошибки 184
- Средних метод 349
- Средняя квадратичная ошибка 172

- Средняя квадратическая ошибка, геометрическое значение 182
- — — суммы двух независимых величин 187
- Средняя ошибка 177
- —, геометрическое значение 184
- Степенная шкала 229
- Схема вычислений для двенадцати ординат 303, 311
- Таблица разностей 264
- Теорема сложения вероятностей 134
- умножения вероятностей 136
- Теория взвешенных измерений 190
- Точность вычислений 76
- измерений 25
- интерполяционных формул 289
- Третья задача общей теории ошибок 120
- Уравнения шкалы 220
- Формула Ньютона 271
- основная теории случайных ошибок 149
- Формулы для приближённых вычислений 51, 373
- Функциональные сетки 231
- шкалы 219
- Функция распределения случайных ошибок 142
- Шкала равномерная 219
- логарифмическая 227
- проективная 229
- степенная 229
- функциональная 219
- Шкалы модуль 220
- уравнение 220
- Экстраполирование 292
- Эмпирические формулы 326
- —, выбор лучшего вида 367

Редактор *Е. Б. Кузнецова*
Техн. редактор *С. С. Гаврилов*
Корректор *Е. А. Белицкая*

Подписано к печати 16/XI 1953 г. Бума-
га 84×108¹/₃₂ 6 бум. л. 19,68 печ. л.
18,40 уч.-изд. л. 37 400 тип. зн. в печ. л.
Т-08265. Тираж 15000 экз. Заказ № 1283.
Цена книги 5 р. 50 к. Переплёт 1 р. 50 к.

16-я типография Союзполиграфпрома
Главиздата Министерства
культуры СССР
Москва, Трёхпрудный пер., 9